

Wasserstein approximation of measures via orthogonal frames: Statistical and computational complexities

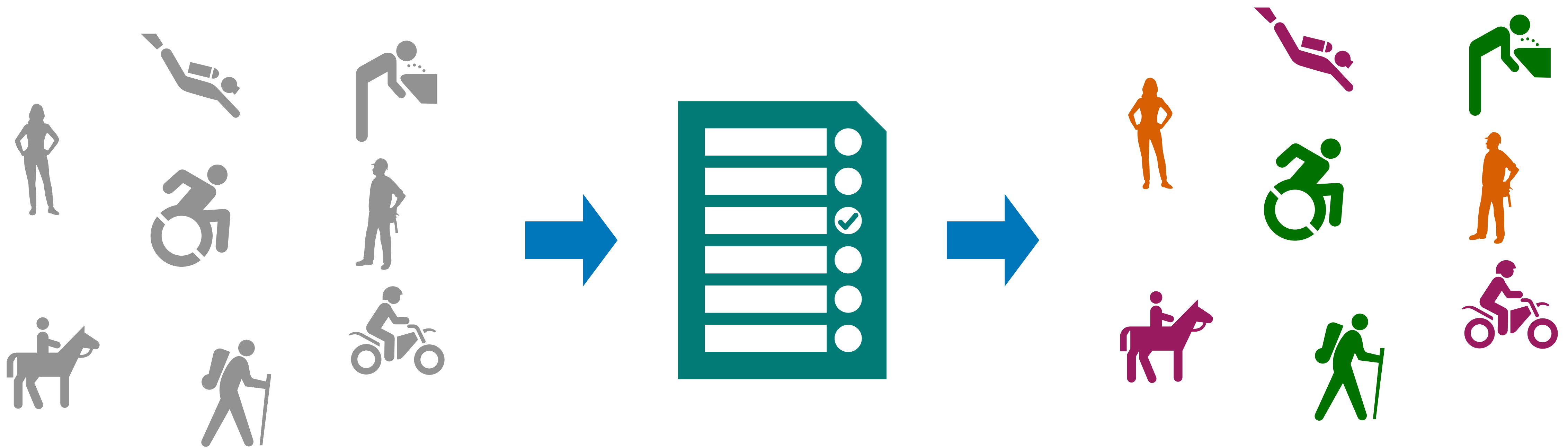
UN Encuentro de Matemáticas (70 años)
August 18th

Daniel López-Castaño
Johns Hopkins University

Joint with Mateo Diaz (U Chicago), Congwei Yang & Nicolás García-Trillos (U Wisconsin-Madison)

Motivation: A problem in Psychometry

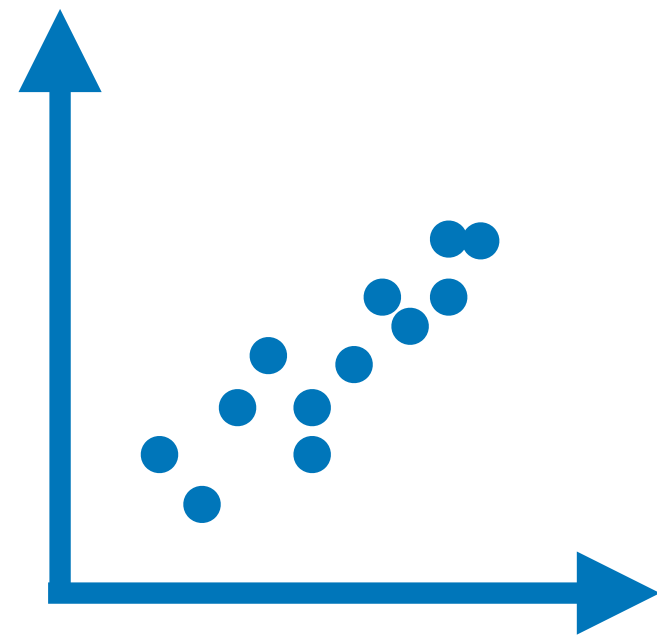
Psychometrics researchers aim to **understand personality types** through survey analysis.



The Mathematical Problem: Given $X \in \mathbb{R}^{d \times n}$, **recover the labels** $y \in \{1, \dots, k\}$

Factor Analysis

Datasets with numerous variables can often be effectively represented using ***fewer latent variables (factors)***



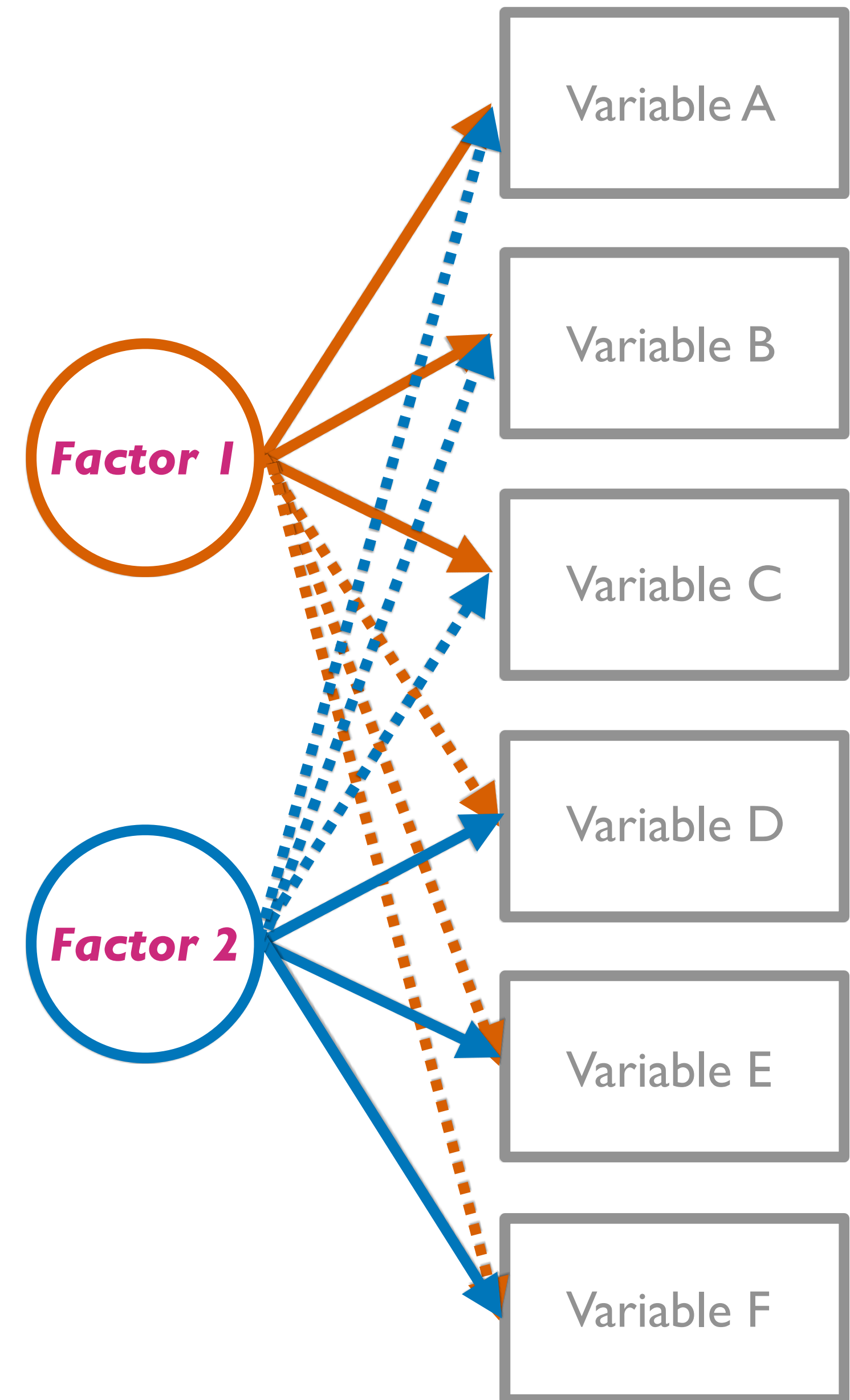
This is often described via a linear model:

$$X = FW^T$$

Data (orange arrow pointing to X)

Factors (pink arrow pointing to F)

Loadings (green arrow pointing to W^T)



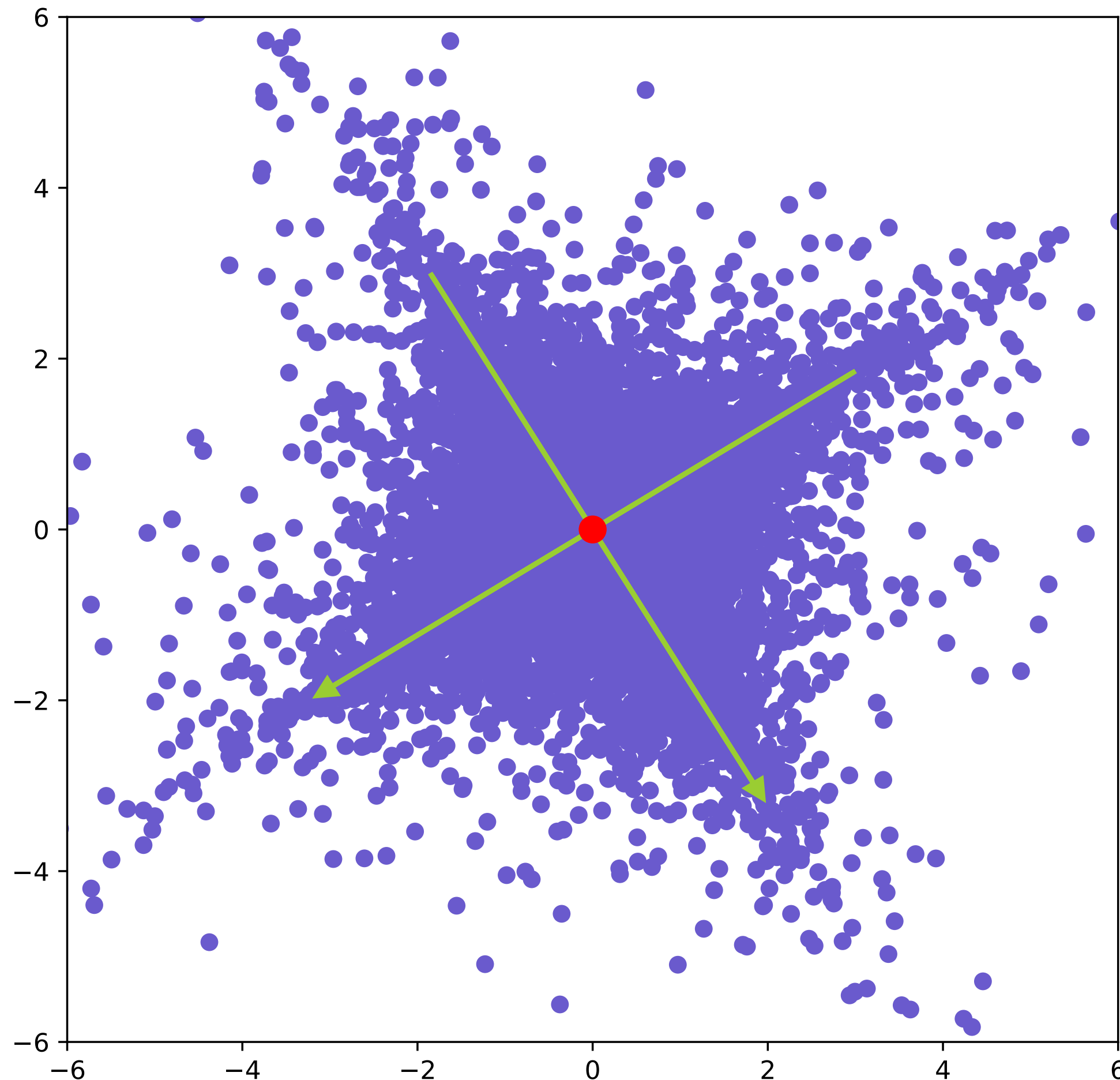
Factor Analysis

Remarkably, in psychometrics, the data often exhibits ***radial structure***

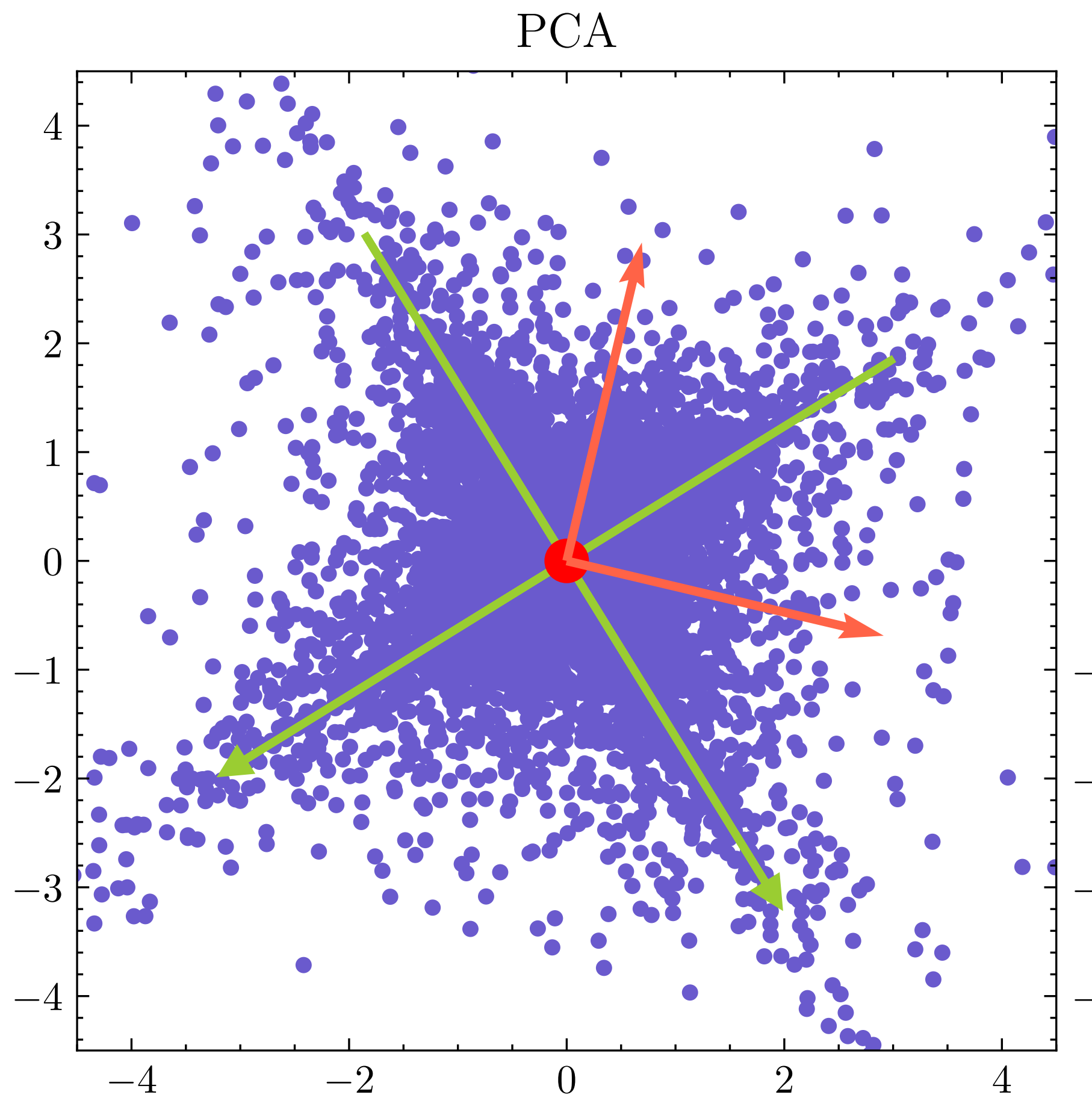
Intuitively, this means that the factorization looks like

$$X_i = Fw_i = [F_1, F_2, \dots, F_k] \begin{pmatrix} 1 \\ \varepsilon_1 \\ \vdots \\ \varepsilon_{k-1} \end{pmatrix}$$

and most information can be captured by **one factor!**



How to recover the radial structure?

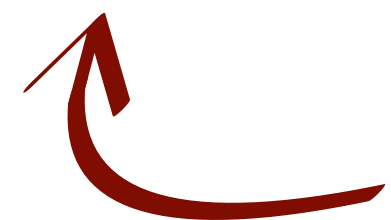


Natural idea: Run PCA.

Problem: The covariance matrix is often close to a multiple of the identity.

$$\mathbb{E} [XX^T] \approx \lambda \text{Id}$$

The **principal components** are essentially **random directions!**



Population covariance is a diagonal matrix

A classical solution: Varimax rotations

Clarify the relationship among **factors**, by **rotating** the orthogonal basis found in PCA

To recover the rotation, i

$$R^* \in \operatorname{argmax}_{R \in O(k, d)} \sum_{i=1}^k \sum_{j=1}^d \left(\sum_{l=1}^p R_{ijl}^4 \right)$$

Intuition: the rotation is determined by fit moment.

PSYCHOMETRIKA—VOL. 23, NO. 3
SEPTEMBER, 1958

THE VARIMAX CRITERION FOR ANALYTIC ROTATION IN FACTOR ANALYSIS*

HENRY F. KAISER
UNIVERSITY OF ILLINOIS

The varimax criterion for analytic rotation in factor analysis

HF Kaiser - Psychometrika, 1958 - Springer

... hoc quality of subjective **rotation** makes uniquely determined **factors** impossible; only **factors** that are ... In contrast, an **analytic** criterion

☆ Save Cite Cited by 11050 Related

[PDF] Factor rotations in factor analyses

H Abdi - Encyclopedia for Research Methods for the Social ..., 2003 - utdallas.edu

... the total subspace after **rotation** is the same as it was before **rotation** (only the partition of ... **rotation** procedures using the loadings of variables **analyzed** with principal component **analysis** ...

☆ Save Cite Cited by 1003 Related articles All 7 versions

Computer program for varimax rotation

HF Kaiser - Educational and psychological measurement, 1958 - Sage

... n tests and r **factors**, the varimax criterion ... 2,..., r are **factors**, a_{ij} is the loading for the jth test on the sth **factor**, and h_i^2 is the communality ...

☆ Save Cite Cited by 524 Related articles All 4 versions

... al normal varimax is appended.

... an analytic criterion for rotation is defined as one ... mathematical conditions beyond the fundamental factor ... theorem, such that a factor matrix is uniquely determined. Historically, the first such criterion was Thurstone's treatment of the principal axes problem [10]: from any arbitrary factor matrix he suggested rotating under the criterion that each factor successively accounts for the maximum variance.

How well does Varimax perform?

In theory (Rohe '22)

Under several assumptions on the distribution from which **factors** can be sampled:

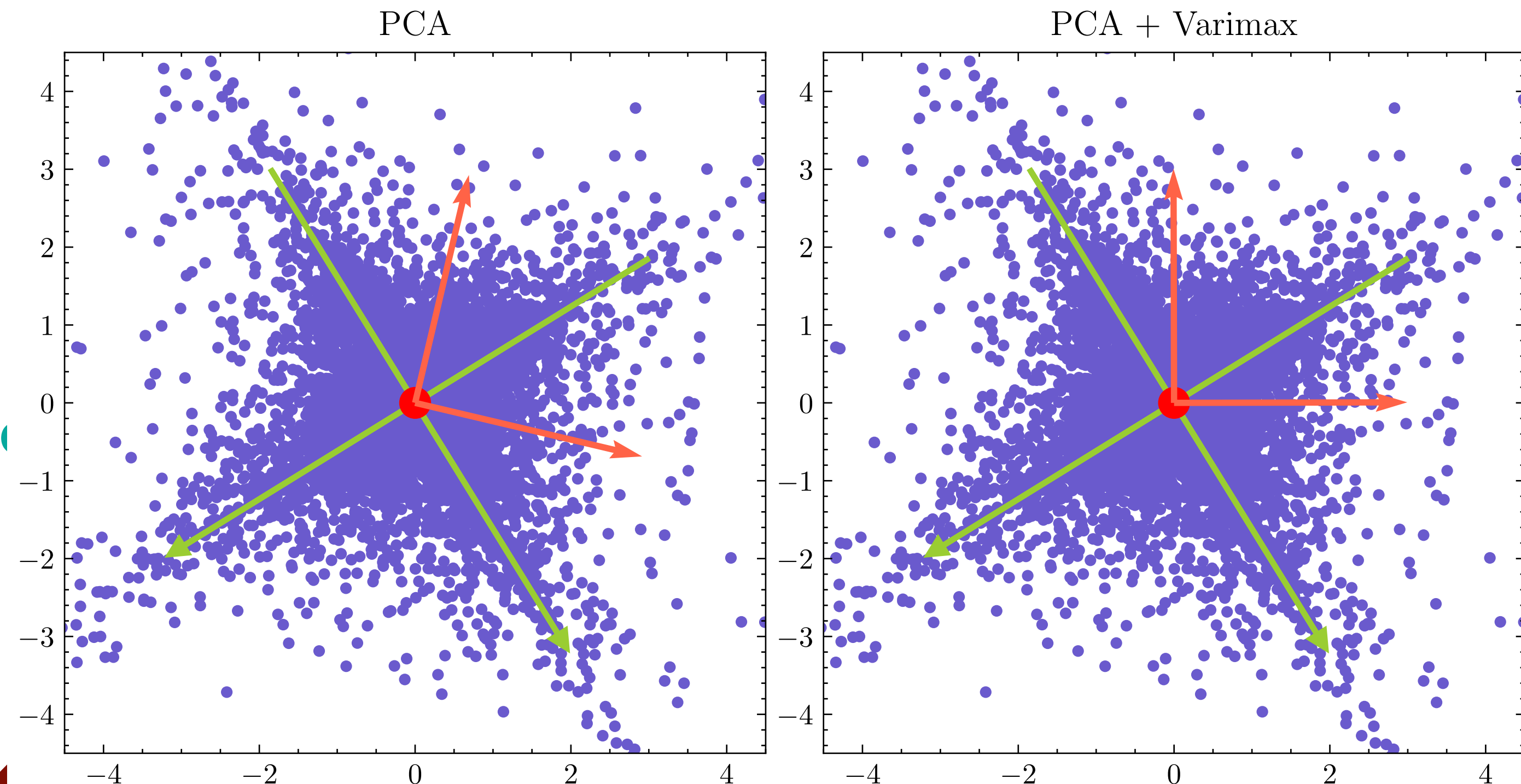
$$\|\hat{X} - XP_n\|_{2 \rightarrow \infty} \leq O_p(\Delta_n^{-0.24} \log^{3.75} n)$$

$n \times$ data matrix

signed-permutation orbit
true factors

factor model

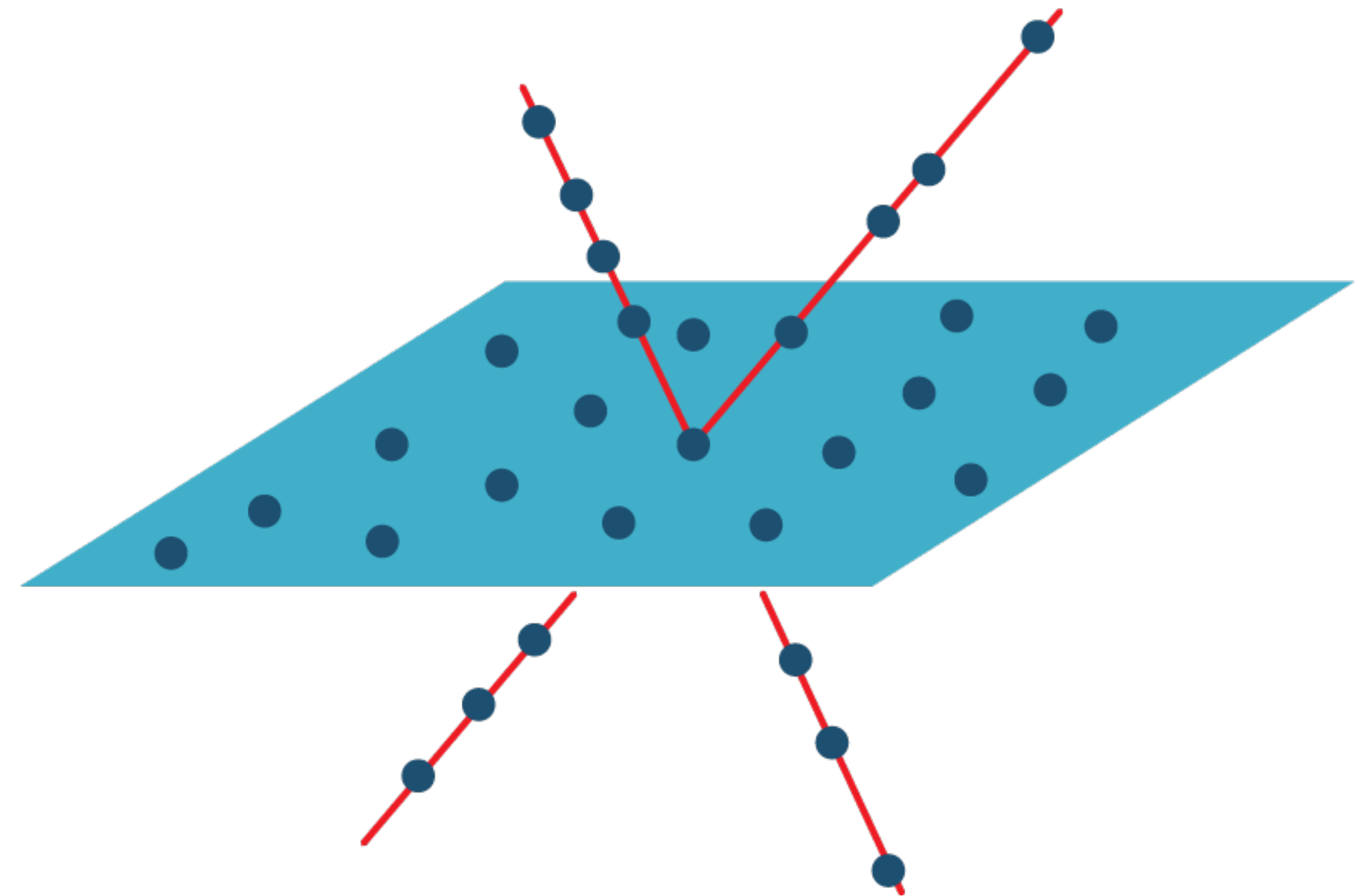
In practice:



Better alternatives?

Subspace clustering aims to recover a **union of subspaces** that capture **most of the data's information**

In particular, our problem is a special case of subspace clustering:
consider k different 1-dimensional subspaces (**factors**)



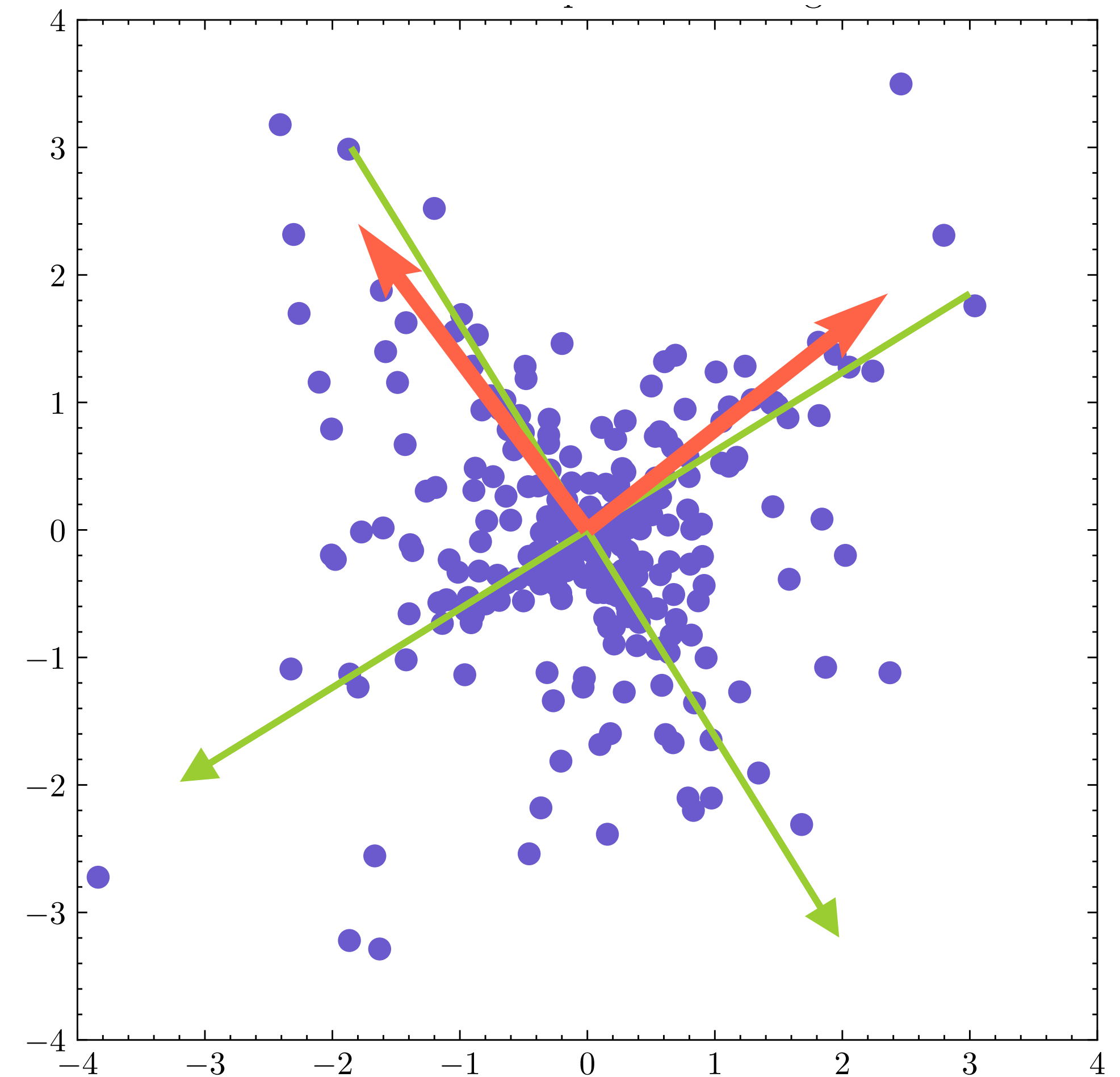
Note: Key applications in computer vision (Vidal, 2011)

Issues with subspace clustering

Issue 1: Most methods **lack robustness**

Issue 2: The most robust solution (Robust subspace clustering, Soltanolkotabi '14) requires **forming an n^2 similarity matrix** and **extracting top- k eigenvalues** on a second stage

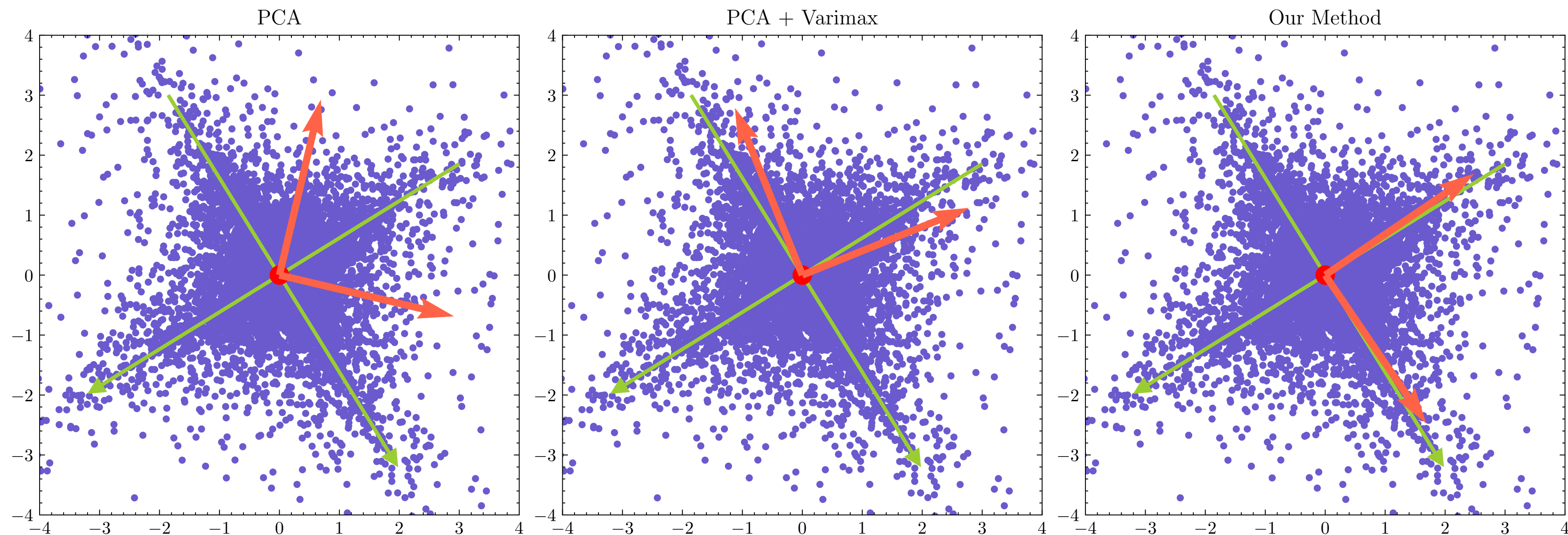
Issue 3: Subspaces are not constrained to be orthogonal to each other, thus it might allow **factors to correlate**



The question of the day

Can we have a scalable method that displays
better statistical performance?

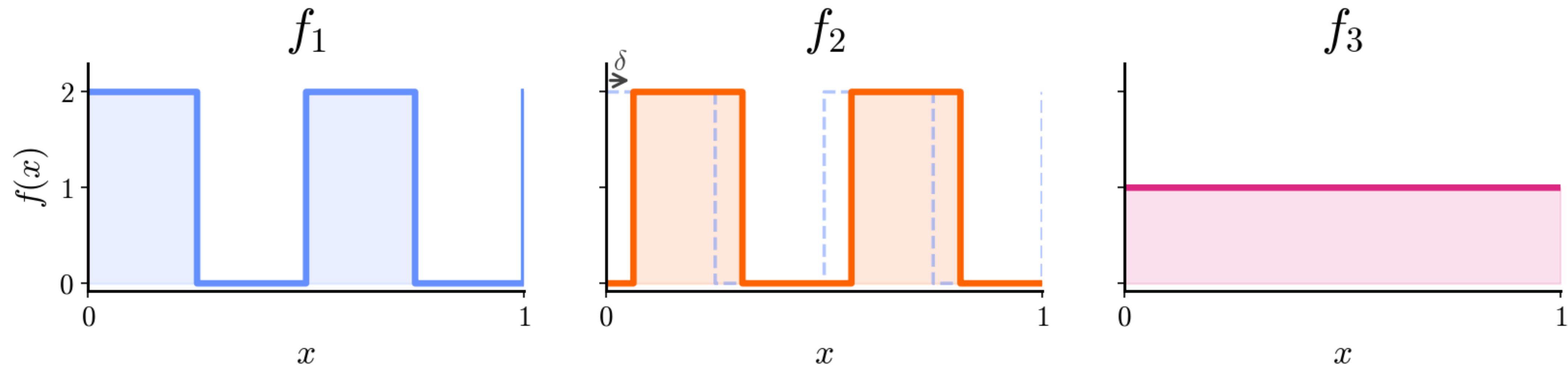
Yes! We introduce **Wasserstein Approx Measure (WAM) w/ manifold opt**



Agenda

- ◆ Motivation and previous work
- ◆ **Short Preliminaries**
- ◆ **Our method**
- ◆ **Statistical Guarantees**
- ◆ **Computational Guarantees**
- ◆ **Numerical Experiments**

Wasserstein Distance

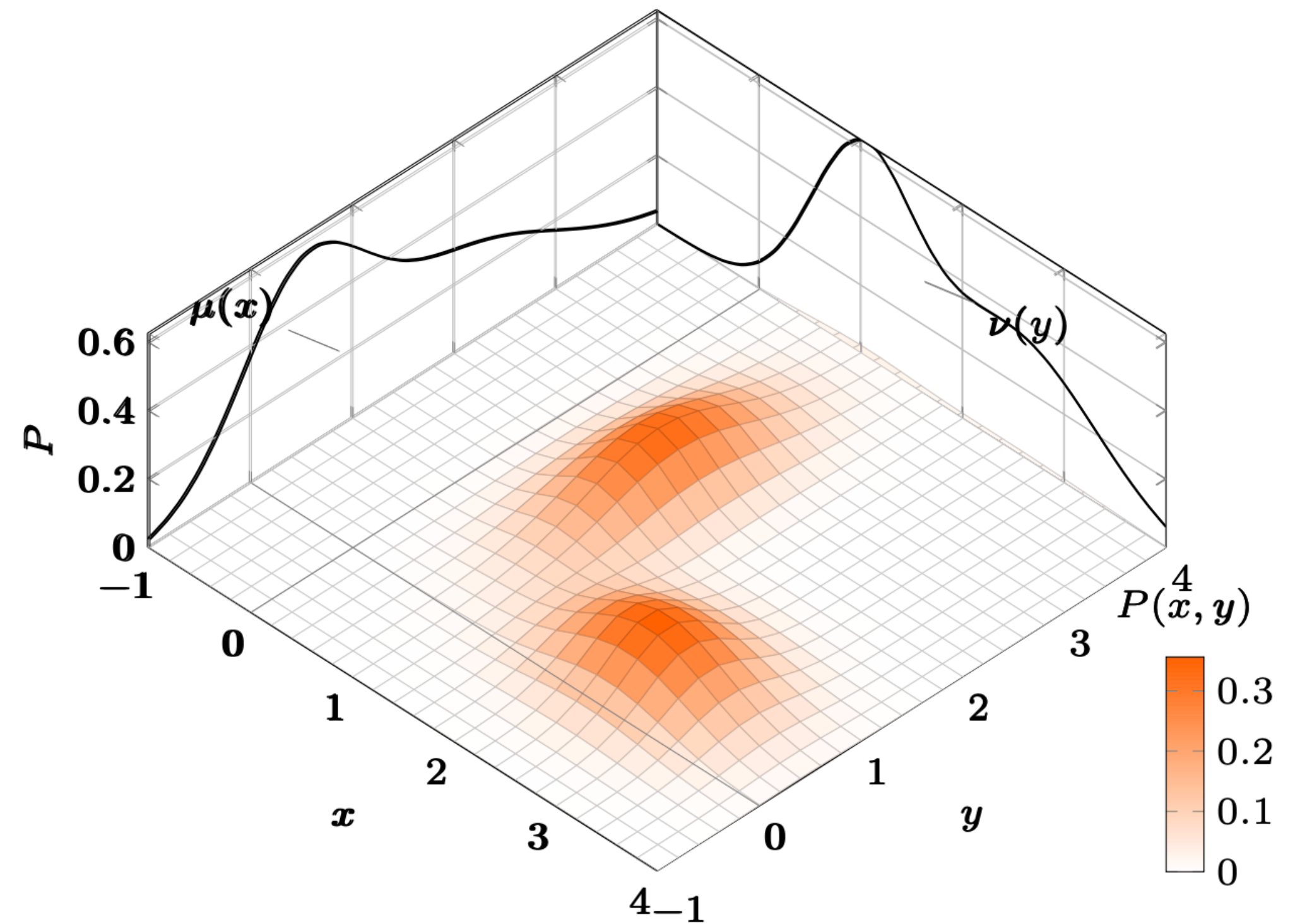
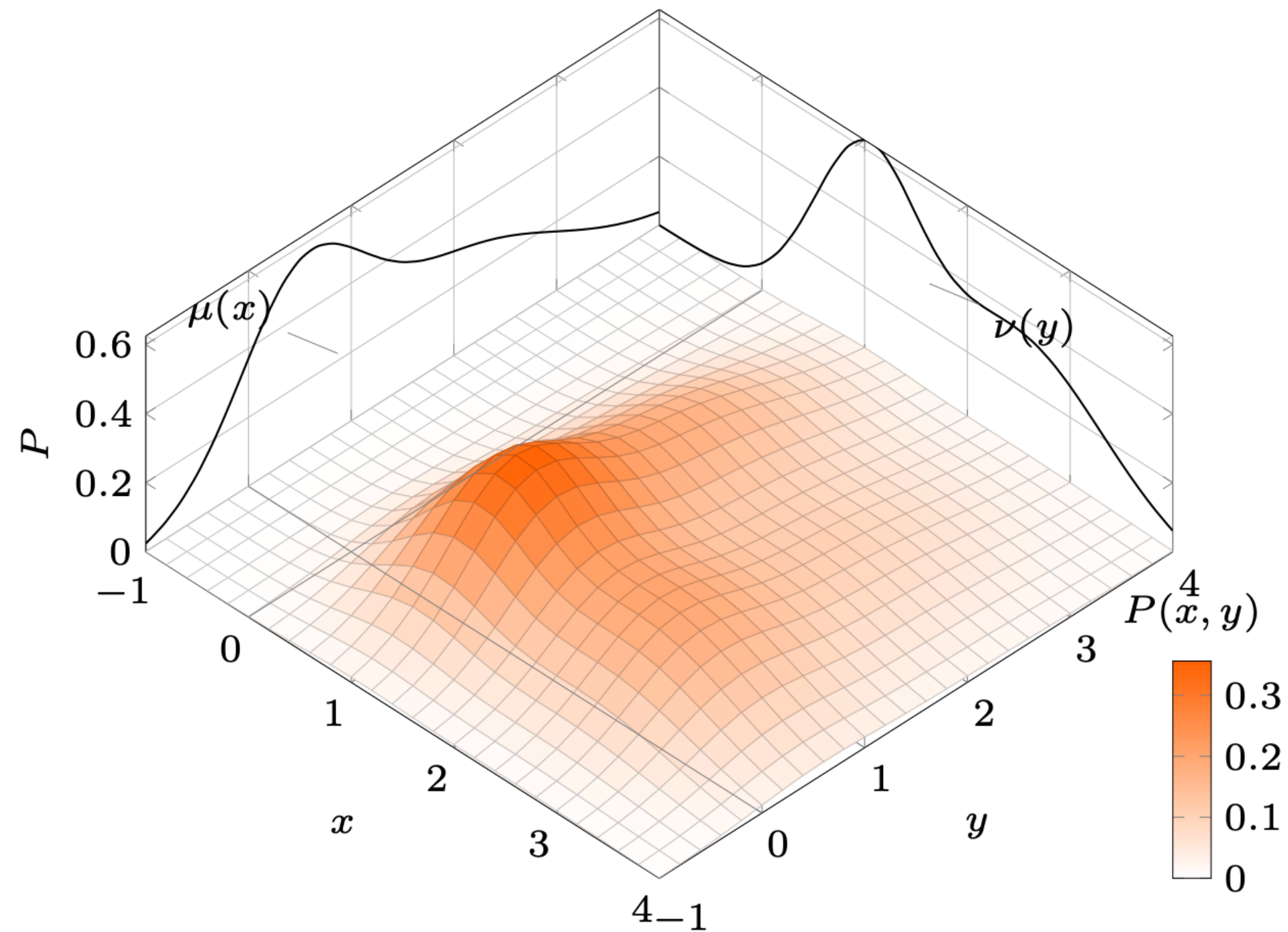


f_2 “should be” closest to f_1 than f_3 , but in L_2 -distance...

$$\|f_1 - f_2\|_2 = \|f_1 - f_3\|_2 = \|f_2 - f_3\|_2 = 1$$

Optimal transport makes it possible to define a **geometry** of a space of probability measures, and thus gives a definition of **distance** in this space.

$$\Pi(\mu, \nu) \stackrel{\text{def}}{=} \{P \in \mathcal{P}(\Omega \times \Omega) \mid \forall A, B \subset \Omega, \\ P(A \times \Omega) = \mu(A), P(\Omega \times B) = \nu(B)\}$$



$$W_2(\mu, \nu) \stackrel{\text{def}}{=} \left(\inf_{P \in \Pi(\mu, \nu)} \iint d(x, y)^2 P(dx, dy) \right)^{1/2}$$

Wasserstein approx measures w/ orthogonal frames

Idea: The data is close to being **supported on a cross**

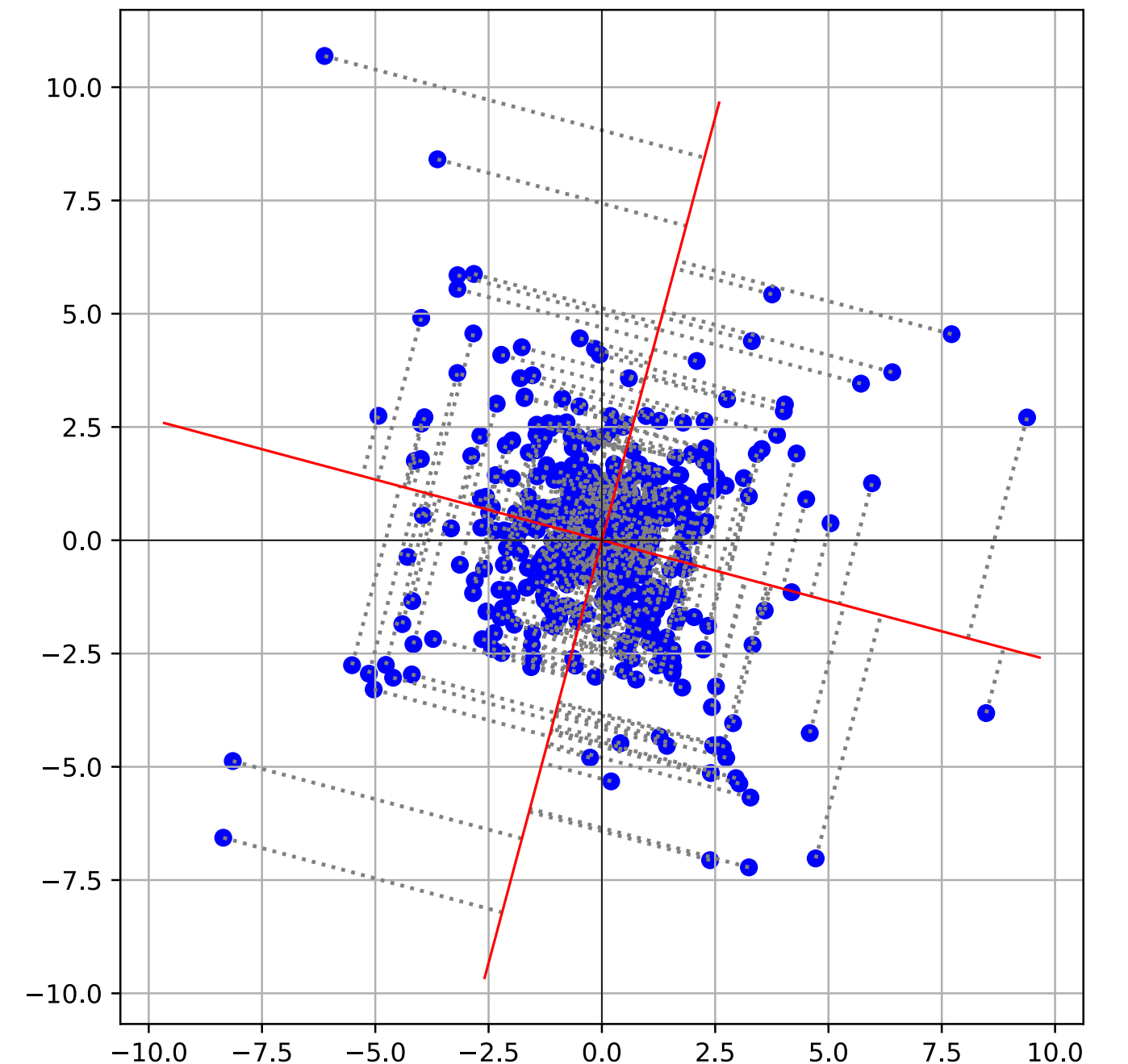
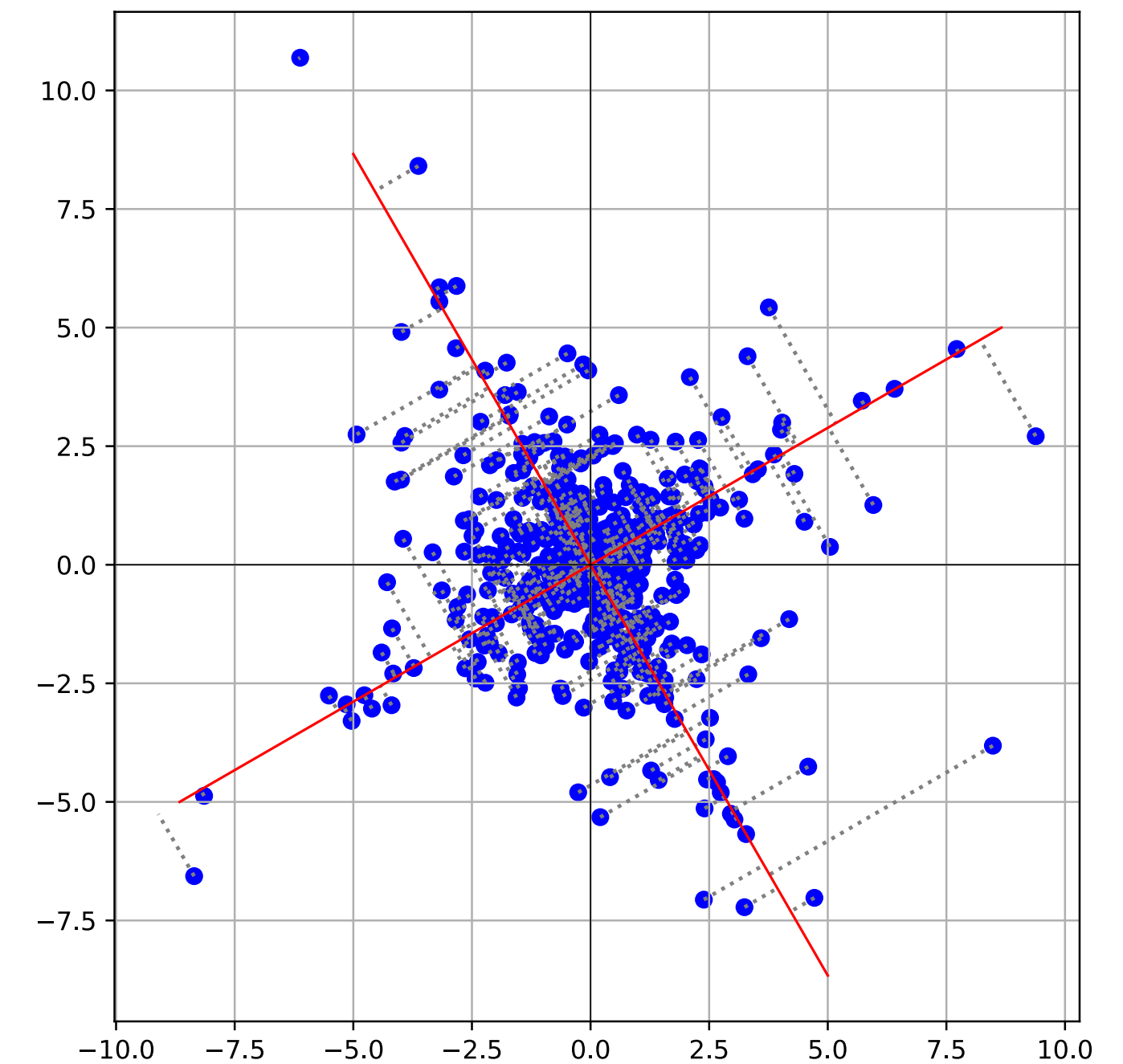
Let's find the closest distribution with that property!

Considering

$\mathcal{P}_\times = \{\text{probability distributions supported on a cross with } k \text{ factors}\},$
we pursue solving

$$\min_{\nu \in \mathcal{P}_\times} W_2^2(\mu, \nu)$$

In principle, this is an **infinite-dimensional** problem,
BUT...



Reduction to finite dimension

Let's consider instead

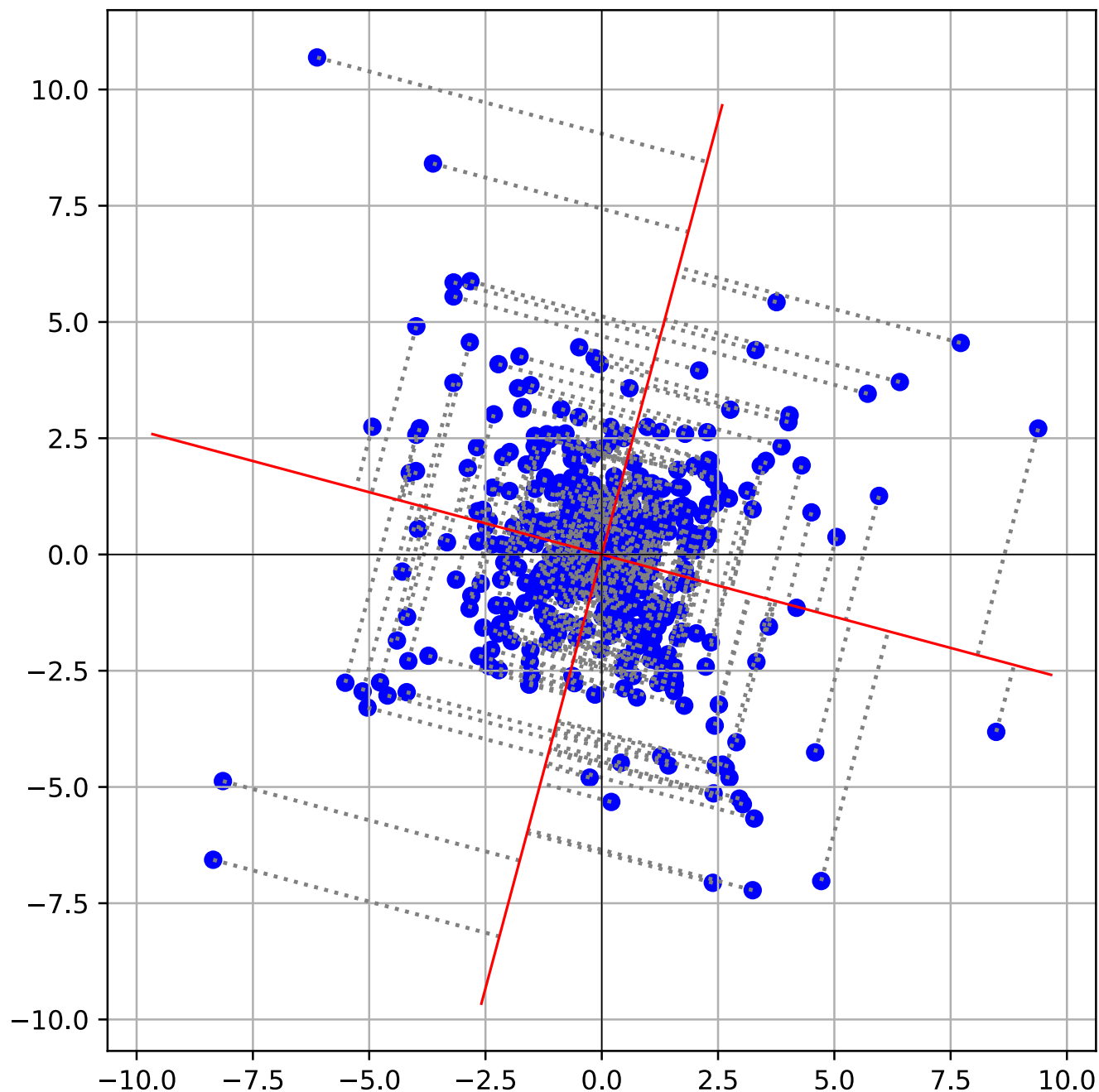
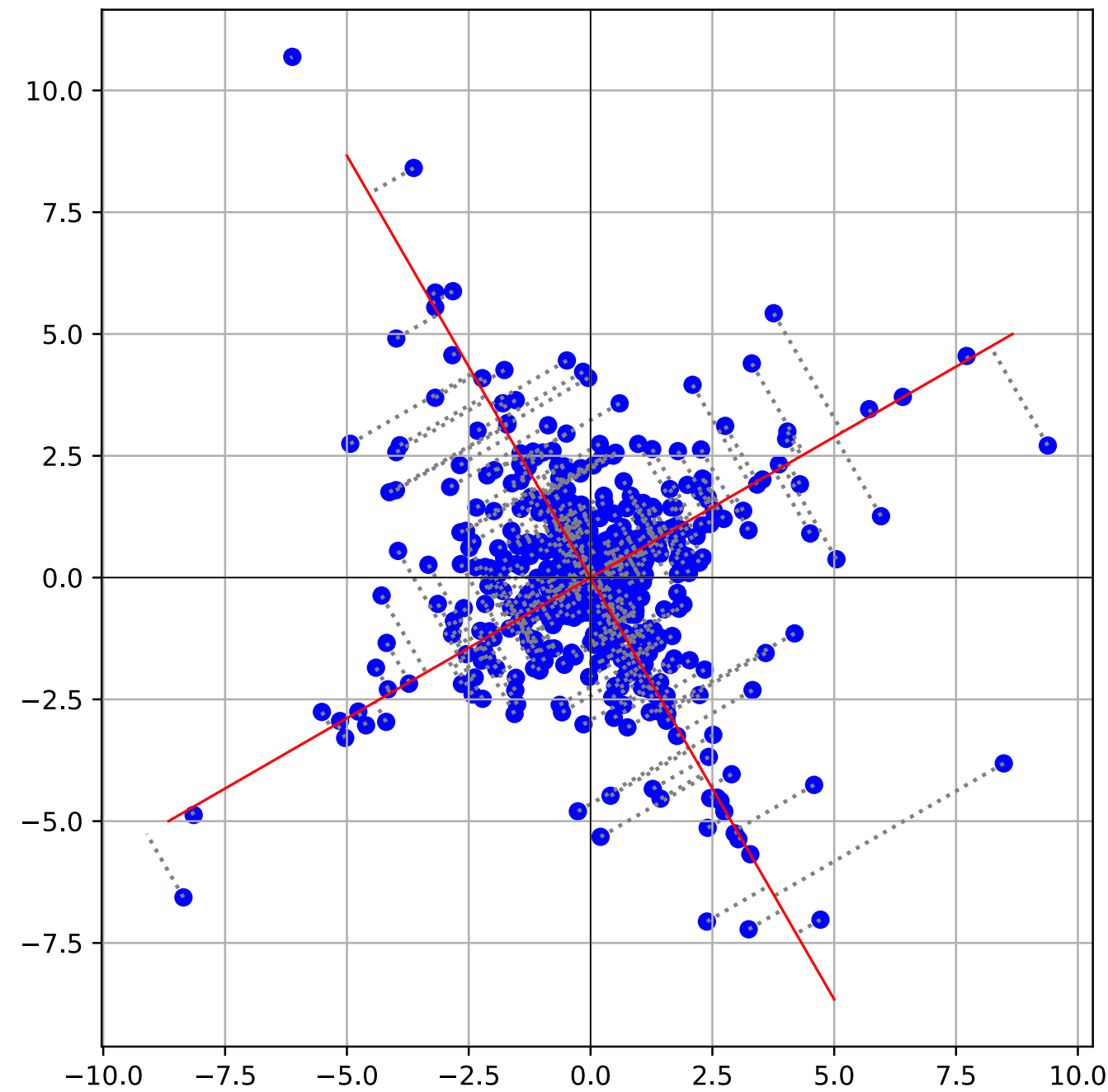
$O(d) := \{U \in \mathbb{R}^{d \times d} : U^T U = I_{d \times d}\}$ and equivalently formulate

$$\min_{U \in O(d)} W_2^2(\mu, \text{proj}_{\times_U} \# \mu)$$

As we've seen, in practice, we are given $x_1, \dots, x_n \sim \mu$ i.i.d., so... we aim to solve

$$\min_{U \in O(d)} \sum_{i=1}^n \|X_i - \text{proj}_{\times_U}(X_i)\|^2$$

Optimize on a manifold?



Interlude: Optimization on Manifolds

Manifold: $O(d) := \{U \in \mathbb{R}^{d \times d} : U^T U = I_{d \times d}\}$

Tangent Space: $T_U \text{O}(d) := \{ U\Omega : \Omega \in \text{Skew}(d) \} = U\text{Skew}(d)$

Retraction: $R_U(X) := (U + X)(I + X^T X)^{-1/2}$

How to optimize without stepping off of the manifold?

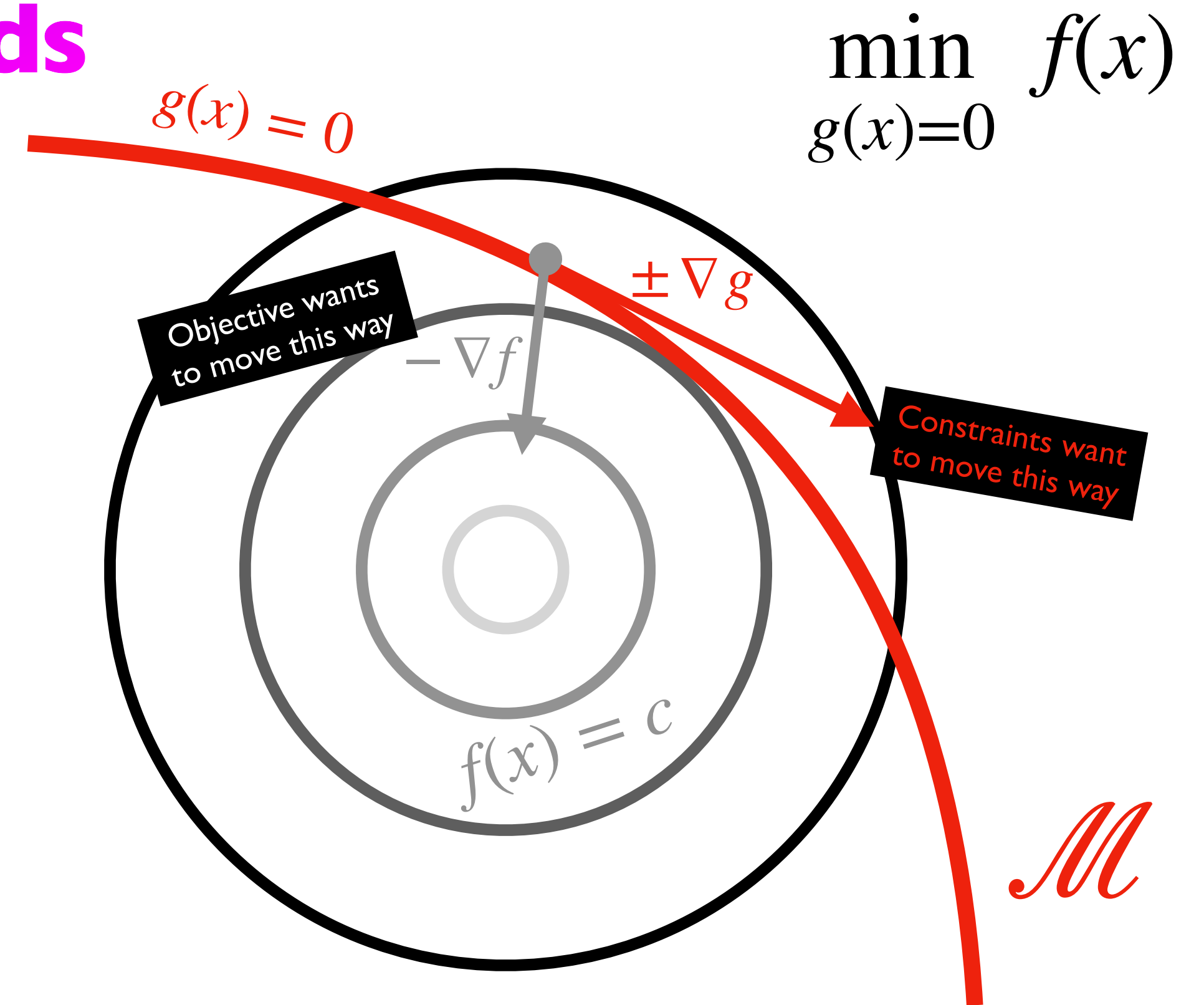
$$x_{k+1} = x_k - \alpha_k \nabla f(x_k) \quad \text{“Euclidean gradient descent”}$$

“Euclidean gradient descent”

$$x_{k+1} = \exp_{x_k}(-\alpha_k \nabla f(x_k))$$

“Manifold gradient descent”

replaced by a weaker
notion: **retraction**



Interlude: Optimization on Manifolds

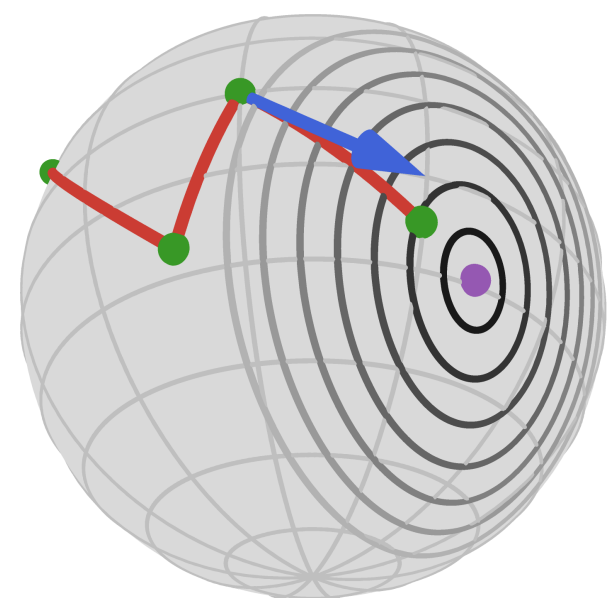
Manifold: $O(d) := \{U \in \mathbb{R}^{d \times d} : U^T U = I_{d \times d}\}$

Tangent Space: $T_U O(d) := \{U\Omega : \Omega \in \text{Skew}(d)\} = U\text{Skew}(d)$

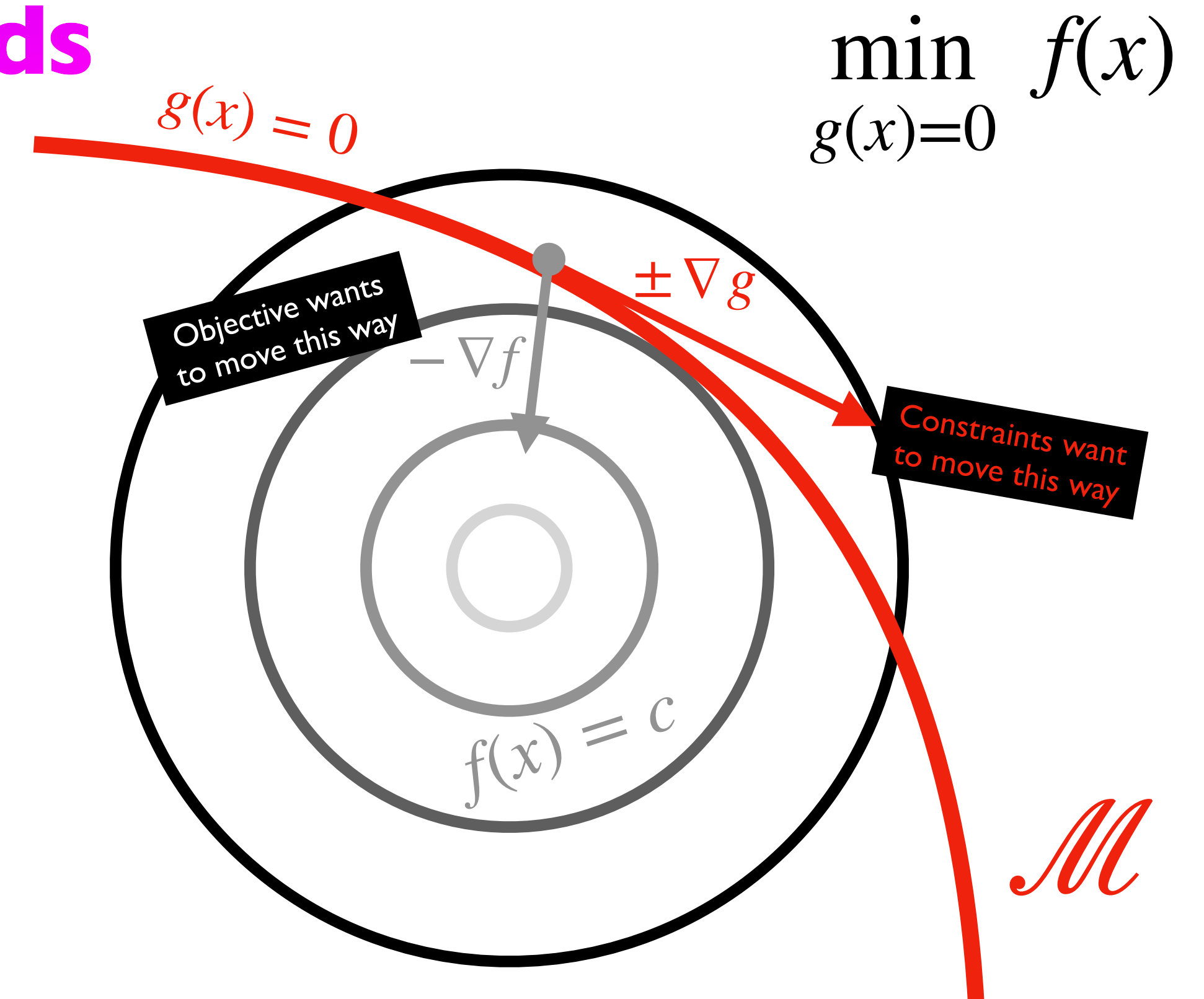
Retraction: $R_U(X) := (U + X)(I + X^T X)^{-1/2}$

Riemannian
Gradient descent!

$$\min_{U \in O(d)} W_2^2(\mu, P_U \# \mu)$$



pymanopt



$$x_{k+1} = x_k - \alpha_k \nabla f(x_k) \quad \text{“Euclidean gradient descent”}$$

walk direction

$$x_{k+1} = \exp_{x_k}(-\alpha_k \nabla f(x_k))$$

“Manifold gradient descent”

replaced by a weaker
notion: **retraction**

Convergence of the empirical minimizer

Theorem I:

Suppose μ is absolutely continuous with finite fourth moment, has sub-exponential squared norms, and satisfies an anti-concentration condition away from the optimal frame axes. Suppose further that frames with sufficiently small population risk $\mathcal{E}(U) := W_2^2(\mu, \text{proj}_{\times_U} \# \mu)$ lie within a geodesically convex neighborhood $\mathcal{N}$ of the population minimizer U^* . Then any empirical minimizer $U_n^* \in \text{argmin}_{U \in \text{O}(d)} \mathcal{E}_n(U)$ satisfies, apart from a distribution-dependent constant, with high probability,

$$d_g(U_n^*, U^*) = \mathcal{O}_p \left(\frac{d^{5/2} \log n}{\sqrt{n}} \right).$$

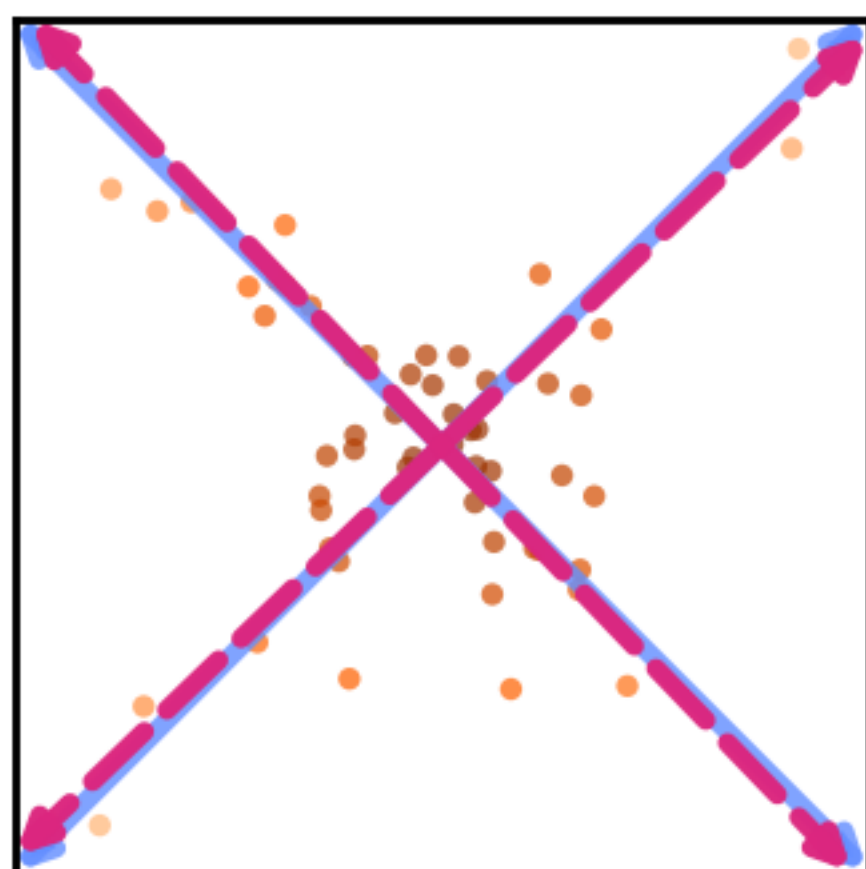
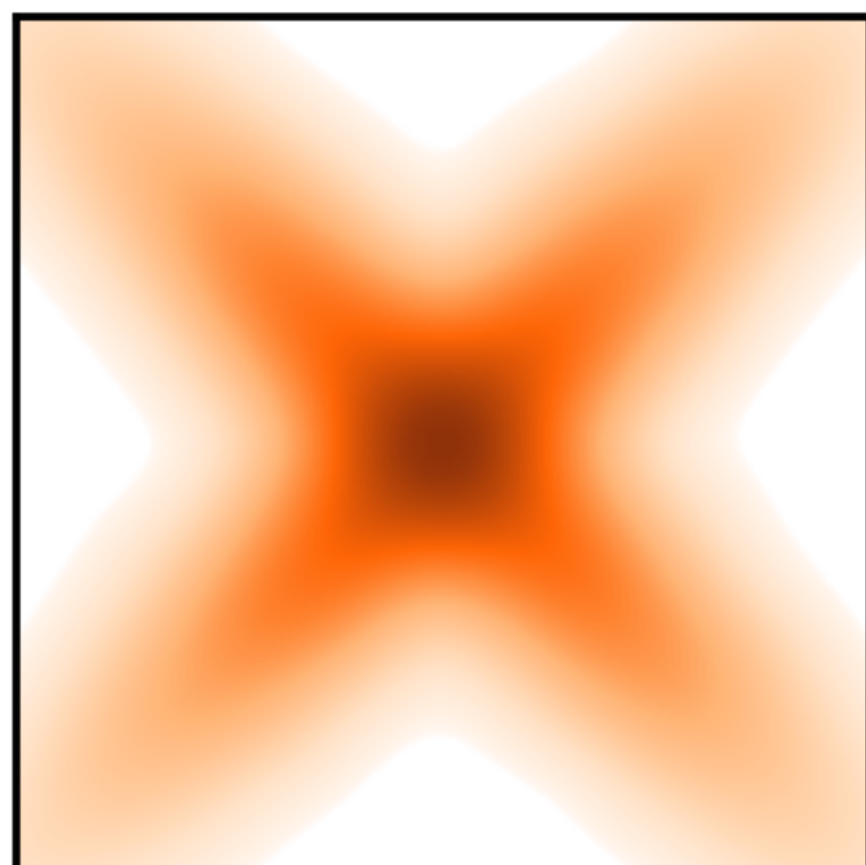
To achieve accuracy ε we need $n = \Omega(d^5 \varepsilon^{-2} \log^2 n)$.

Standard M-estimation steps:

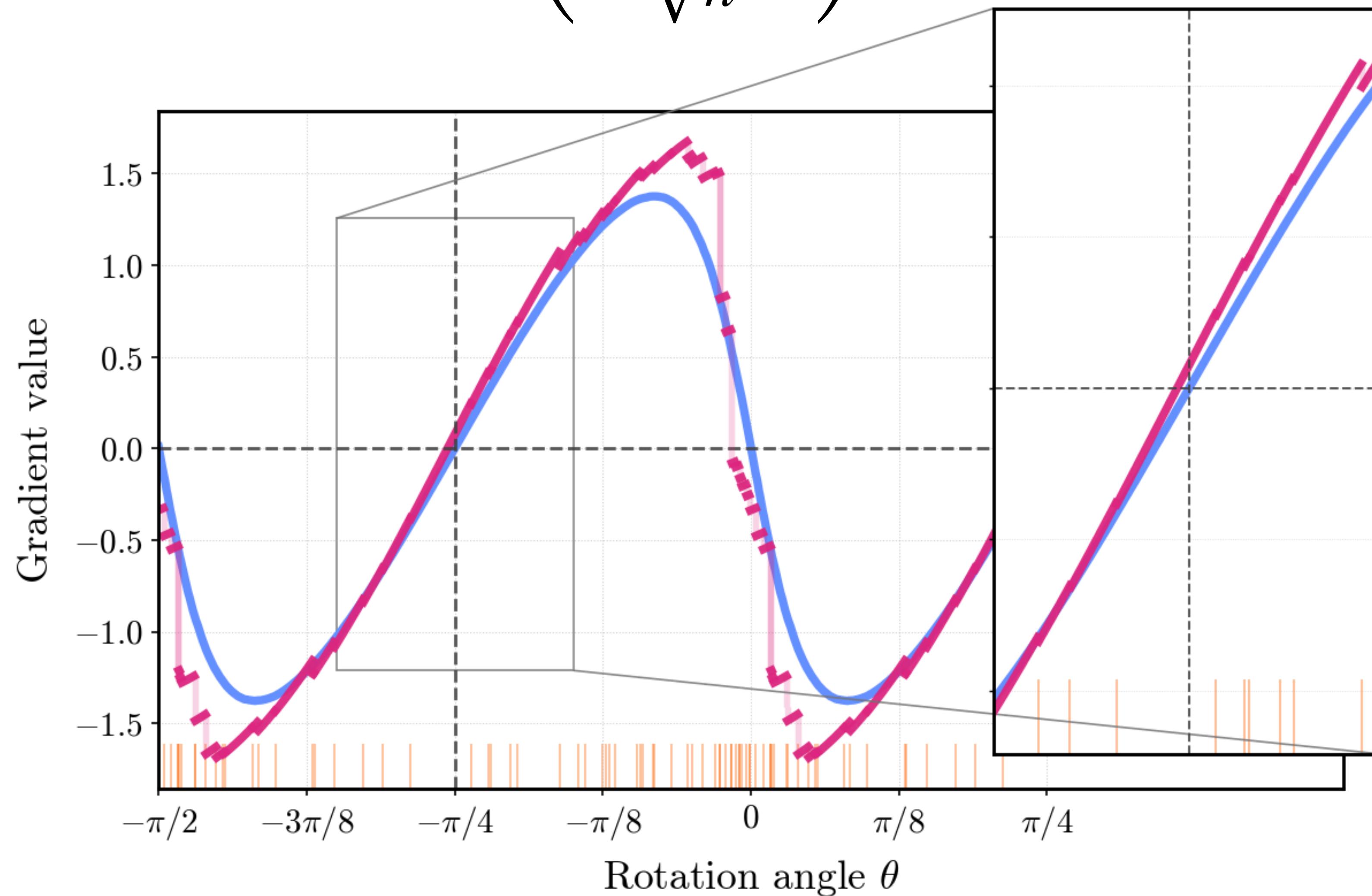
- ✱ CLT gives you $O_p(n^{-1/2})$ control of the gradient error at the true minimizer;
- ✱ Lipschitz continuity of the gradient lets you extend that pointwise bound to a whole neighborhood, so it also applies at the random, data-dependent empirical minimizer;
- ✱ strong convexity converts small gradient into small parameter error.



$$d_g(U_n^*, U^*) = \mathcal{O}_p \left(\frac{d^{5/2} \log n}{\sqrt{n}} \right).$$



Population minimizer
Empirical minimizer



Population gradient Empirical subgradient Axis switches

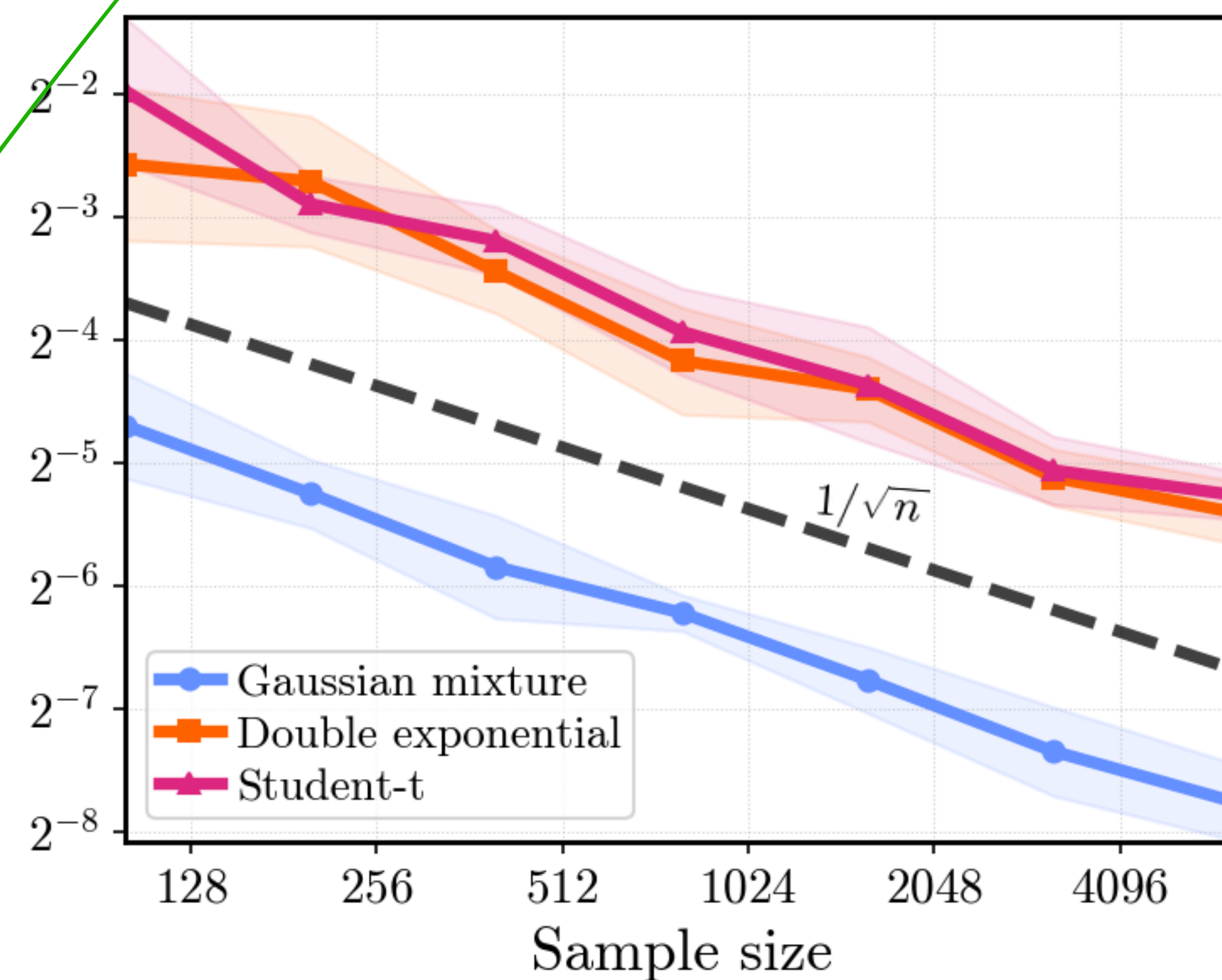
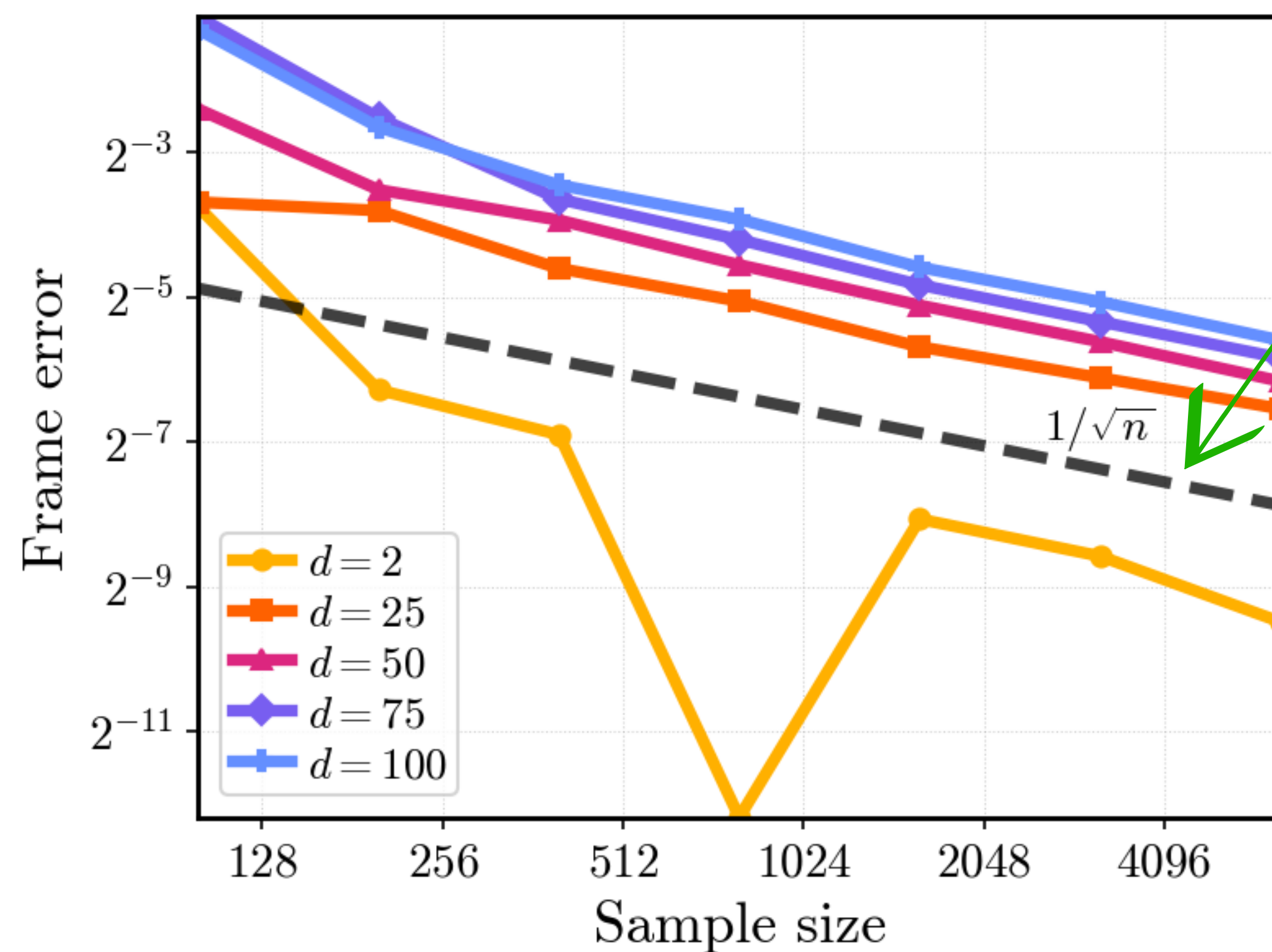
$\nabla^R \mathcal{E}(U)$

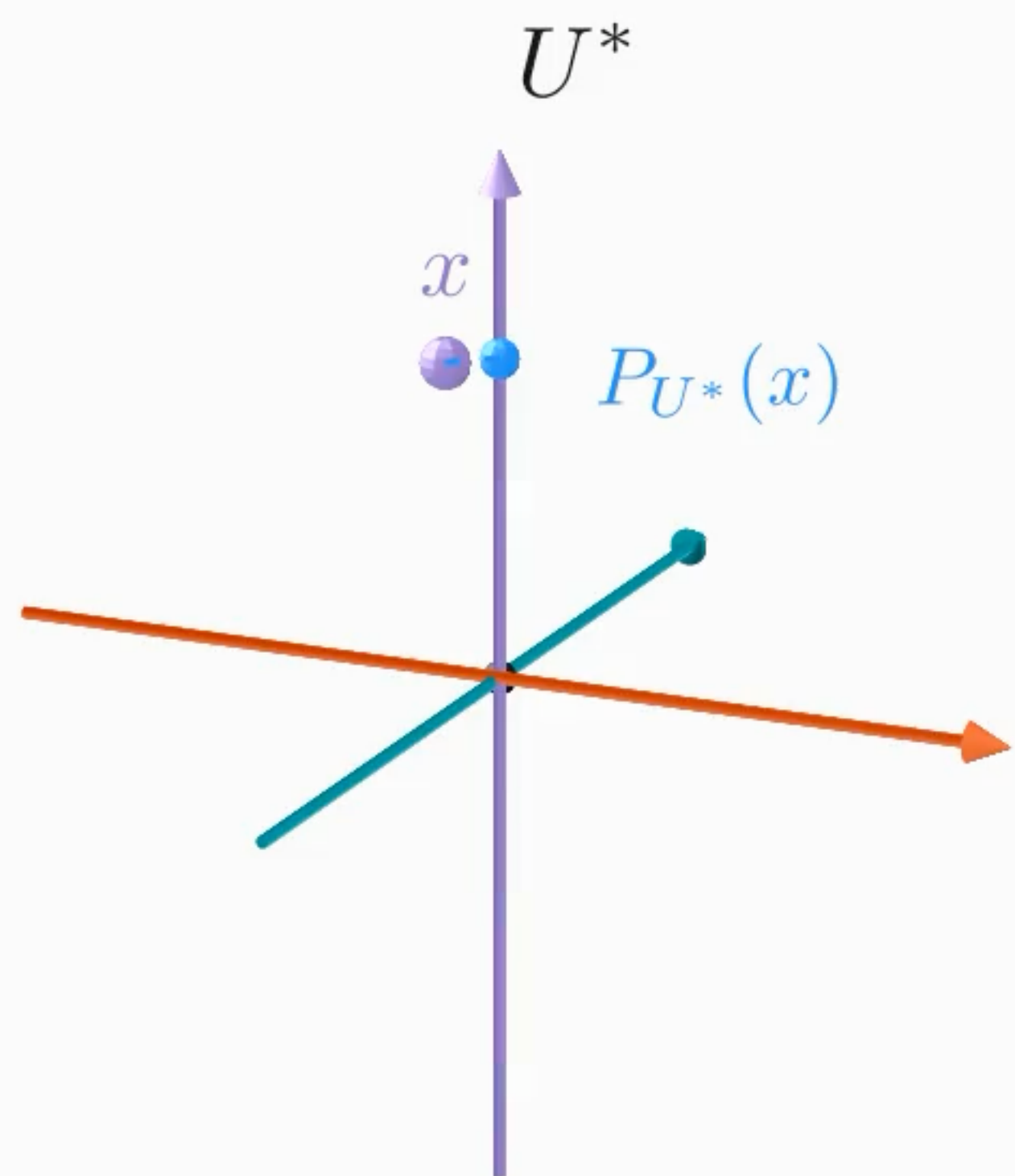
$\partial^R \mathcal{E}_n(U)$

How?

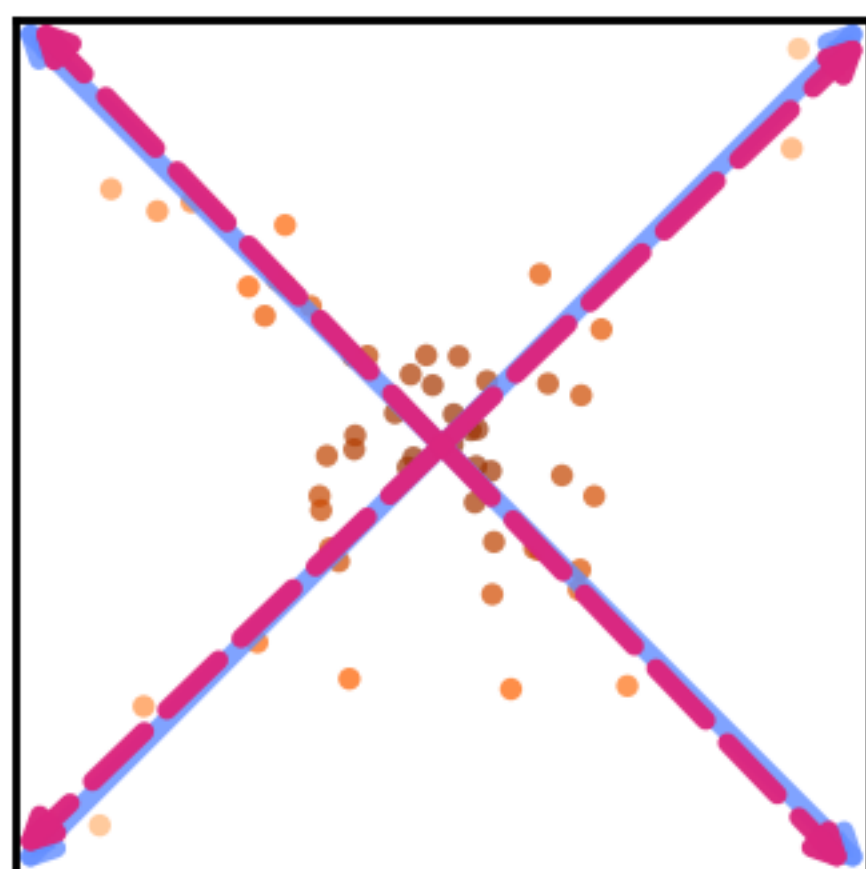
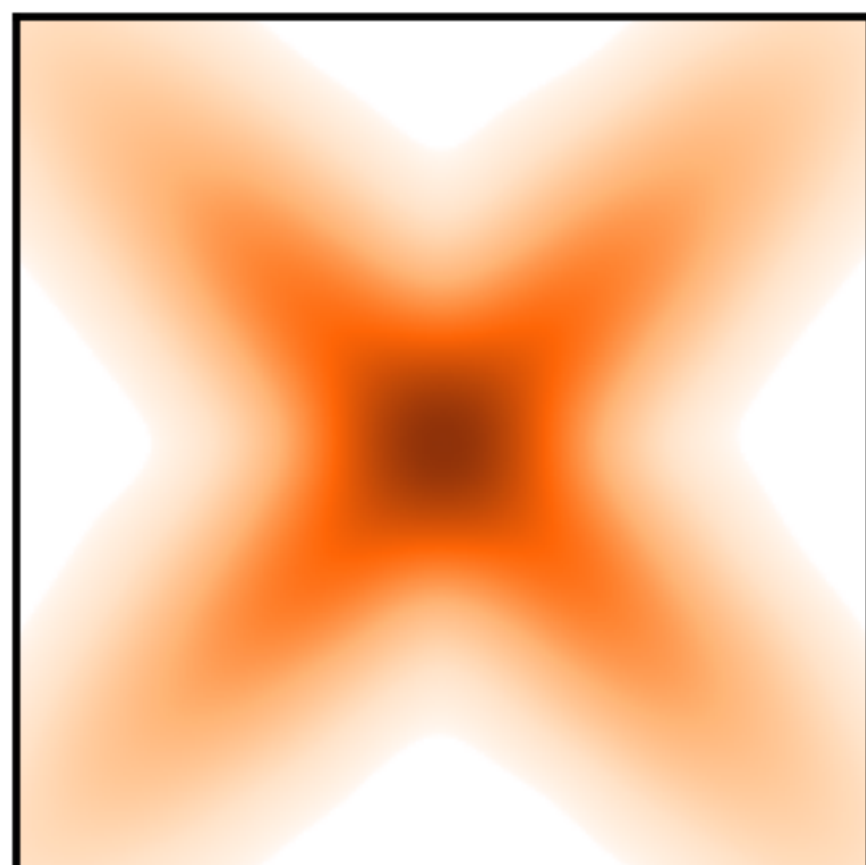
$$d_g(U_n^*, U^*) = \mathcal{O}_p \left(\frac{d^{5/2} \log n}{\sqrt{n}} \right).$$

No curse of dimensionality!

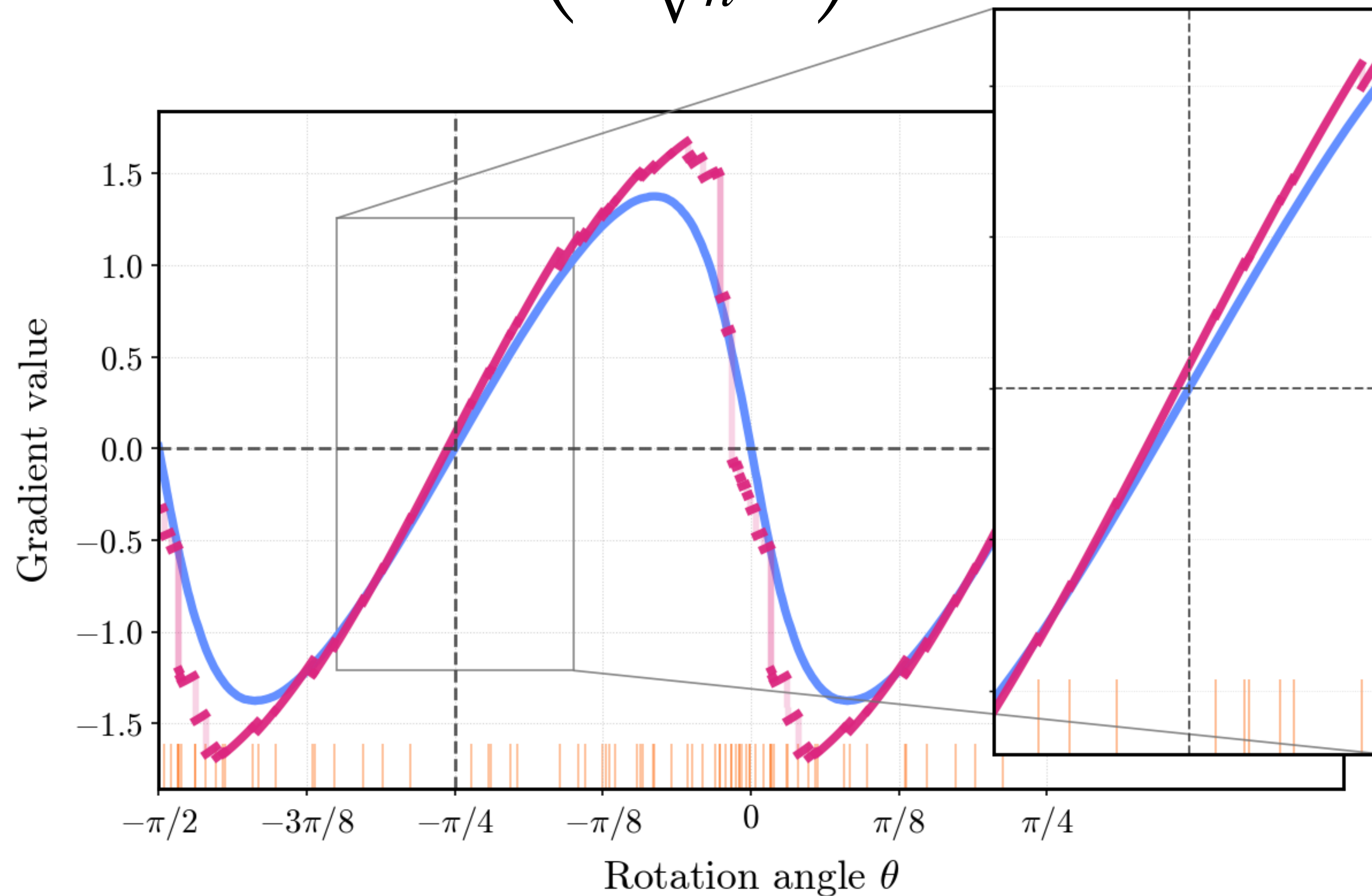




$$d_g(U_n^*, U^*) = \mathcal{O}_p \left(\frac{d^{5/2} \log n}{\sqrt{n}} \right).$$



Population minimizer
Empirical minimizer



Population gradient Empirical subgradient Axis switches

$$\nabla^R \mathcal{E}(U)$$

$$\partial^R \mathcal{E}_n(U)$$

Standard M-estimation steps:

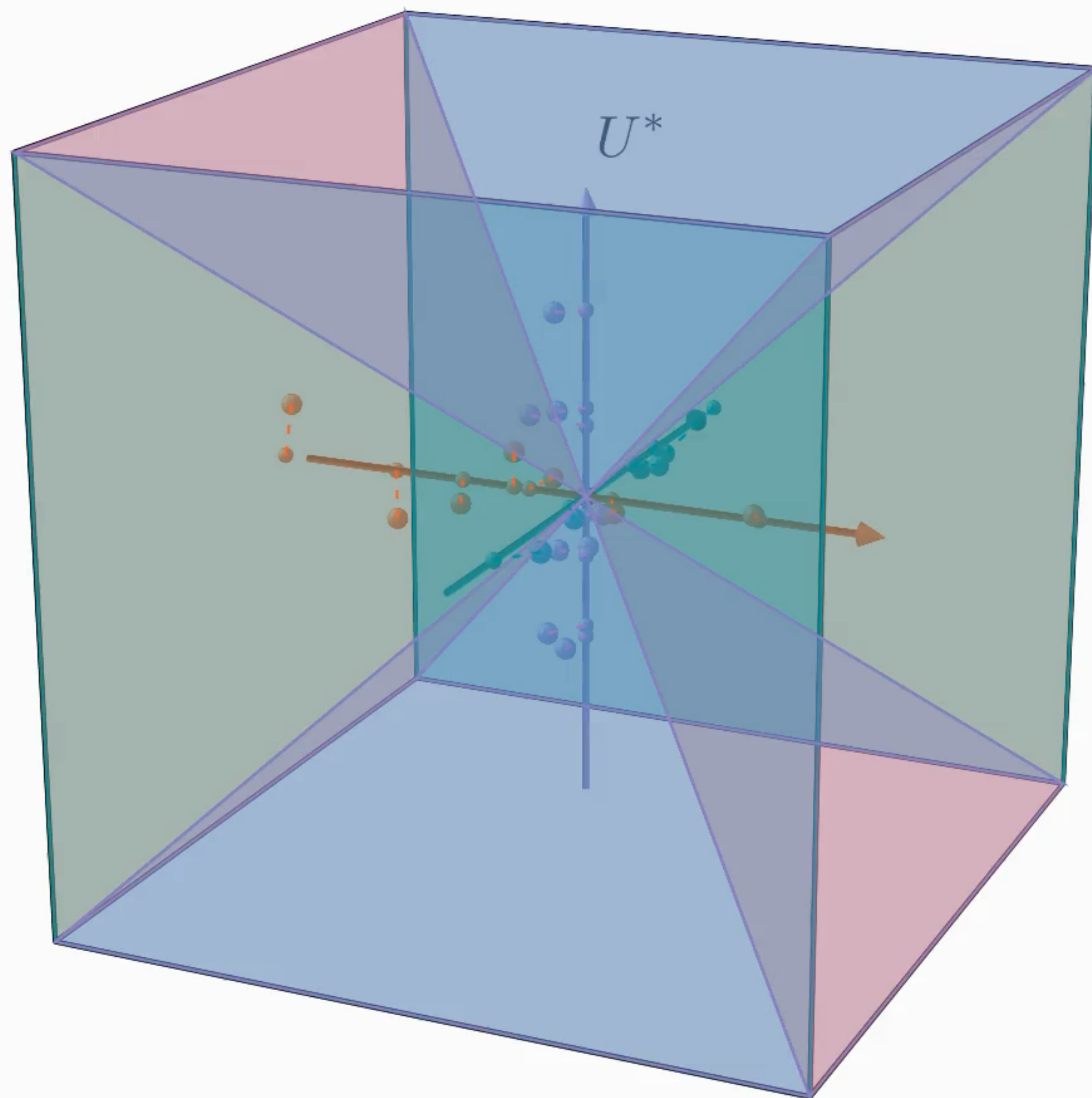
- * CLT gives you $O_p(n^{-1/2})$ control of the gradient error at the true minimizer;
- * Lipschitz continuity of the gradient lets you extend that pointwise bound to a whole neighborhood, so it also applies at the random, data-dependent empirical minimizer;
- * strong convexity converts small gradient into small parameter error.

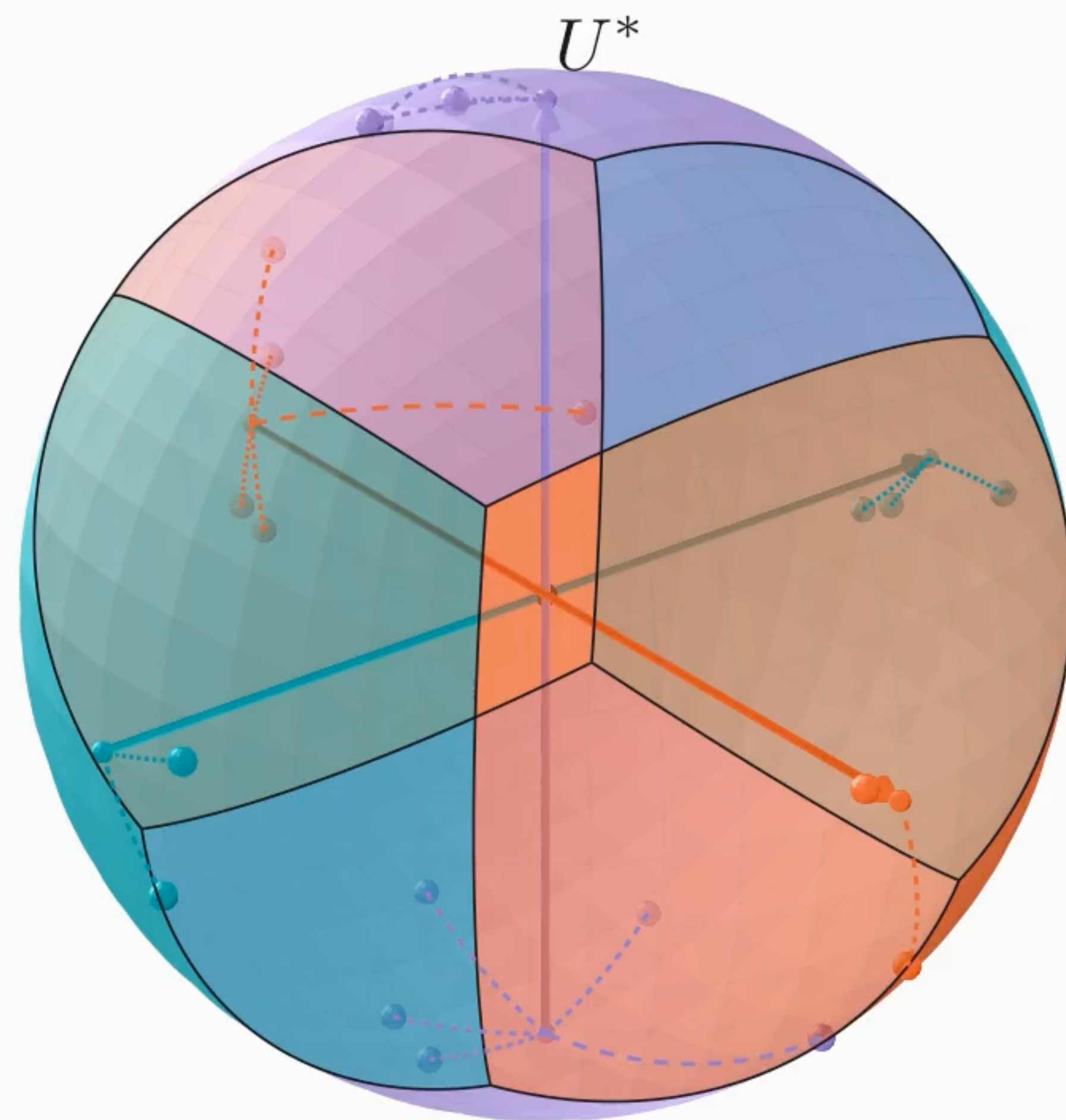


quasi-Lipschitzness

$$\|G_n(U_1) - G_n(U_2)\|_{\text{Fr}} \leq \underbrace{4\|x\|_2^2 \cdot \|U_1 - U_2\|_{\text{Fr}}}_{\text{ordinary Lipschitz term}} + \underbrace{Cd(d-1)\frac{\omega_{d-1}}{\omega_d}\frac{\log n}{\sqrt{n}}}_{\text{resolution error (switching)}} \text{ w/ } \|U_1 - U_2\|_{\text{Fr}} \leq 1/\sqrt{n}$$

$$\underbrace{Cd(d-1)\frac{\omega_{d-1}}{\omega_d}\frac{\log n}{\sqrt{n}}}_{\text{resolution error}} = \underbrace{\frac{1}{n}}_{\text{average over } n} \times \underbrace{N_{1/\sqrt{n}}(U_1, U_2)}_{\text{\# samples that switch axis}} \times \underbrace{O(\log n)}_{\substack{\text{max jump per switch,} \\ \|x_i\|^2 \lesssim \log n}}$$





Theorem [Switching count concentration]:

For any $U_1, U_2 \in \mathcal{N}$ such that $\|U_1 - U_2\|_{\text{Fr}} \leq 1/\sqrt{n}$, with high probability,

$$N_{1/\sqrt{n}}(U_1, U_2) = d(d-1) \left(C \frac{\omega_{d-1}}{\omega_d} \sqrt{n} + \dots \right)$$

where ω_d denotes the surface area of the unit sphere in $\mathbb{R}^d$.

$$\underbrace{C d(d-1) \frac{\omega_{d-1}}{\omega_d} \frac{\log n}{\sqrt{n}}}_{\text{resolution error}} = \underbrace{\frac{1}{n}}_{\text{average over } n} \times \underbrace{N_{1/\sqrt{n}}(U_1, U_2)}_{\substack{\text{\# samples that switch axis}}} \times \underbrace{O(\log n)}_{\substack{\text{max jump per switch,} \\ \|x_i\|^2 \lesssim \log n}}$$

$$d_g(U_n^*, U^*) \leq \frac{1}{m} \left\| \nabla_R \mathcal{E}(U_n^*) \right\|_F \quad (\text{geodesic strong convexity})$$

$$(\text{uniformity argument}) \leq \frac{1}{m} \sup_{U \in \mathcal{N}} \left\| \nabla_R \mathcal{E}(U) - G_n(U) \right\|_F$$

$$(\varepsilon\text{-net argument}) \leq \frac{1}{m} \left(\max_{U_k} \left\| \nabla_R \mathcal{E}(U_k) - G_n(U_k) \right\|_F + \max_{U_k} \sup_{\|U - U_k\| \leq \delta} \left\| G_n(U) - G_n(U_k) \right\|_F \right)$$

$$(\text{pointwise CLT + quasi-L}) \leq \frac{1}{m} \left(C_1 d \sqrt{\frac{Cd(d-1)\log n}{n}} + 4 \|x\|_{L^2}^2 \delta + C_2 d(d-1) \frac{\omega_{d-1}}{\omega_d} \frac{\log n}{\sqrt{n}} \right)$$

$$(\text{absorbing constants}) \leq \frac{1}{m} \left(C_1 d^2 \sqrt{\frac{\log n}{n}} + C_2 d(d-1) \frac{\omega_{d-1}}{\omega_d} \frac{\log n}{\sqrt{n}} \right)$$

$$= O_p \left(\frac{d^{5/2} \log n}{\sqrt{n}} \right)$$

Computational guarantees

Theorem II:

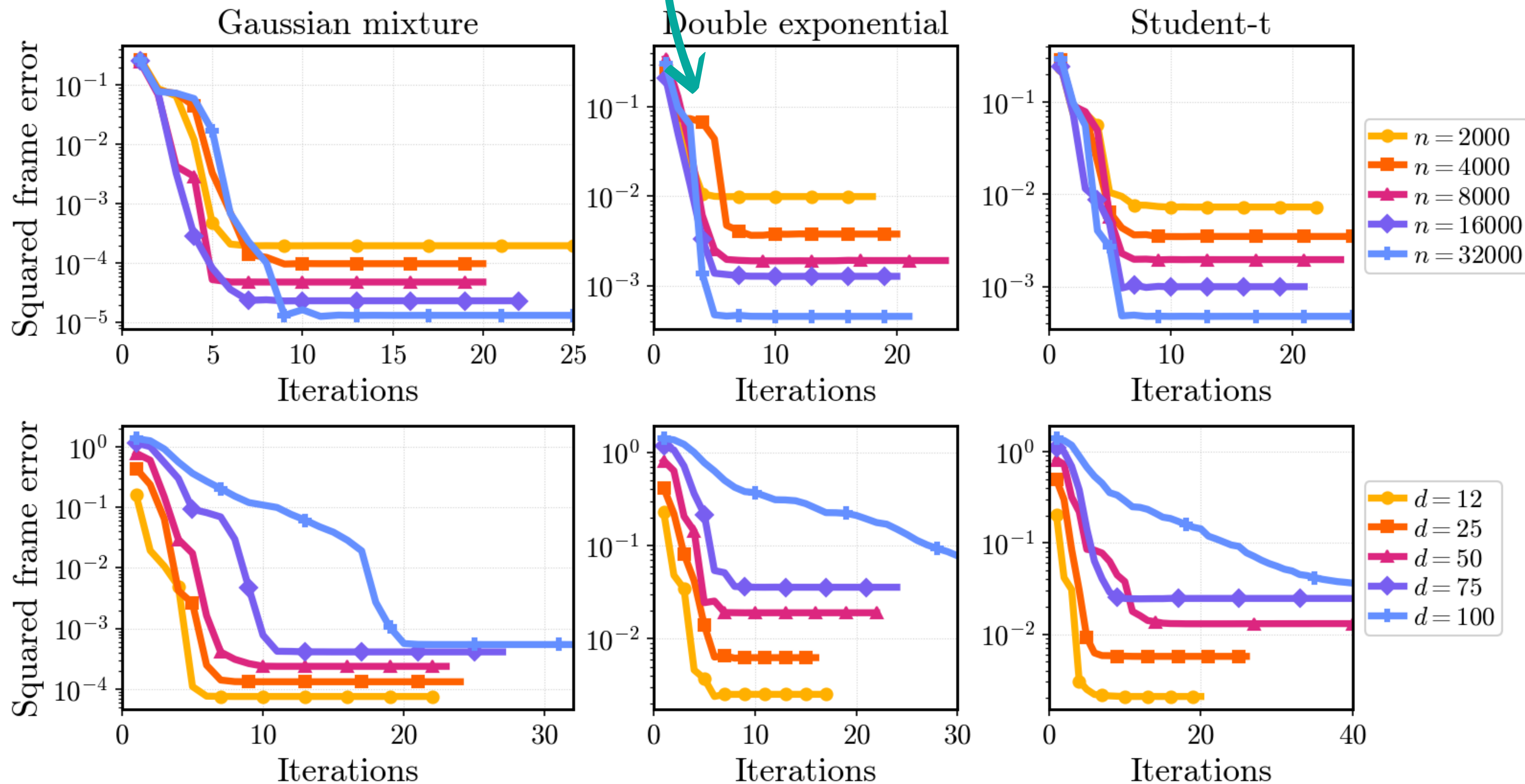
Suppose μ is absolutely continuous with finite fourth moment and that the population risk $\mathcal{E}(U)$ is geodesically strongly convex on $\mathcal{N}$. The Riemannian online-SGD algorithm on $O(d)$ with step size $\eta \in (0, 2m/L^2)$, run for $T = O(\log n)$ iterations, initialized at $U_0 \in \mathcal{N}$, produces a frame U_T satisfying

$$\mathbb{E} \left[d_g(U_T, U^*)^2 \right] = O\left(\frac{\log(n)^3}{n} \right).$$

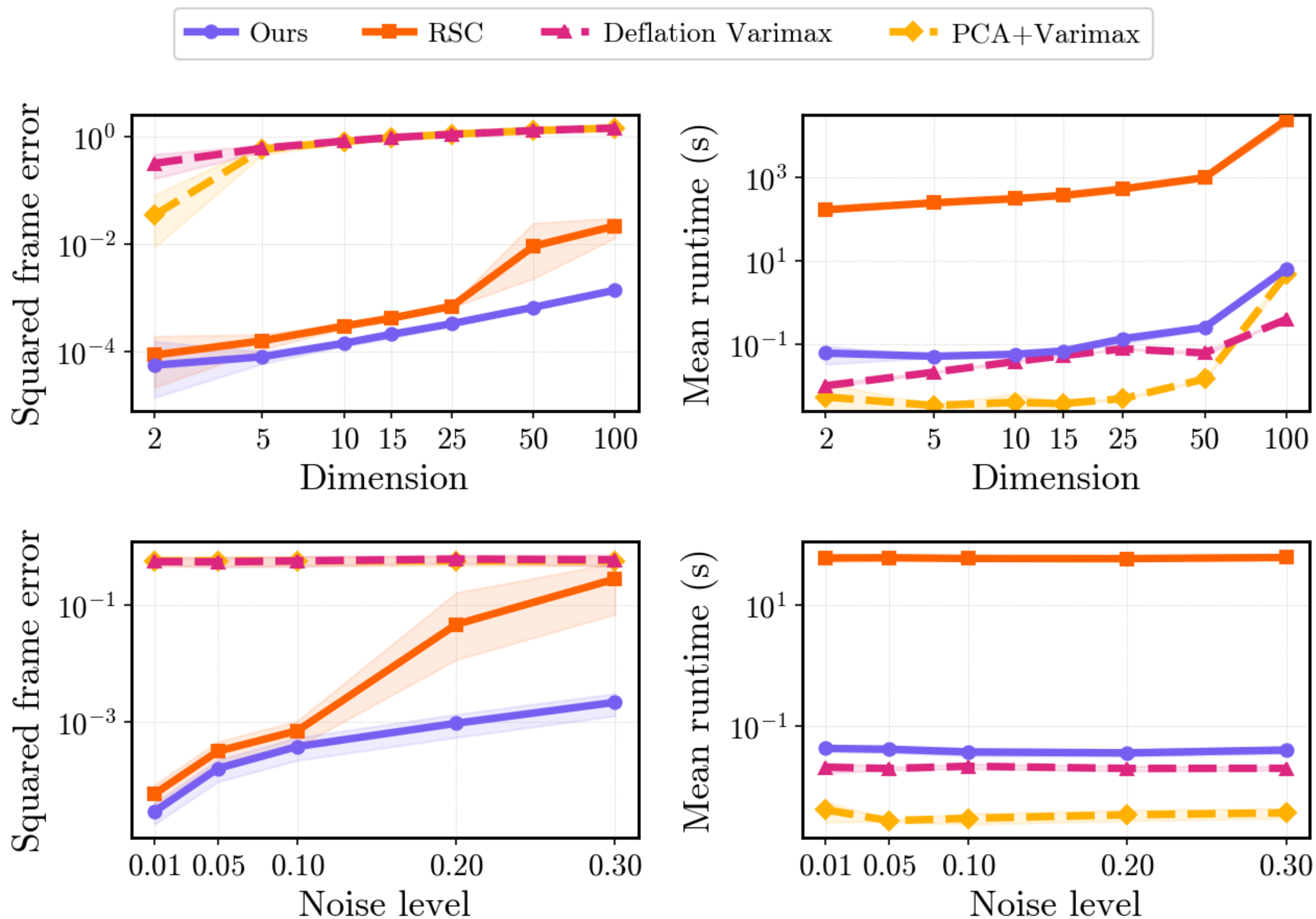
Under the additional condition that the squared norms are sub-exponential, the same rate holds with high probability, with an exponentially small failure probability. The constants m and L quantify the strong convexity and smoothness of the population objective $\mathcal{E}$ near U^* .

$$\mathbb{E} \left[d_g(U_T, U^*)^2 \right] = O\left(\frac{\log(n)^3}{n} \right).$$

Steep initial slope

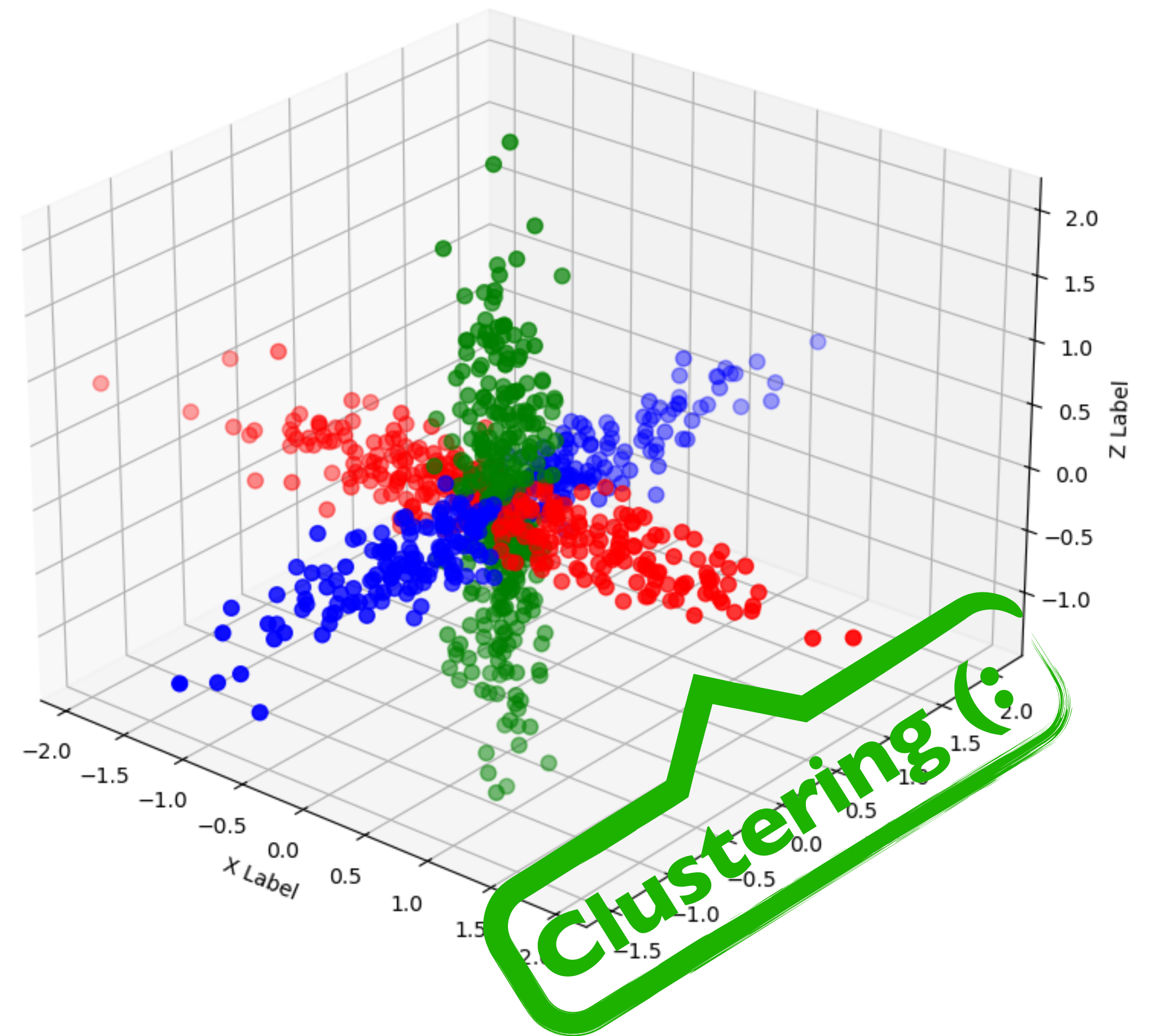


Compared to other methods...

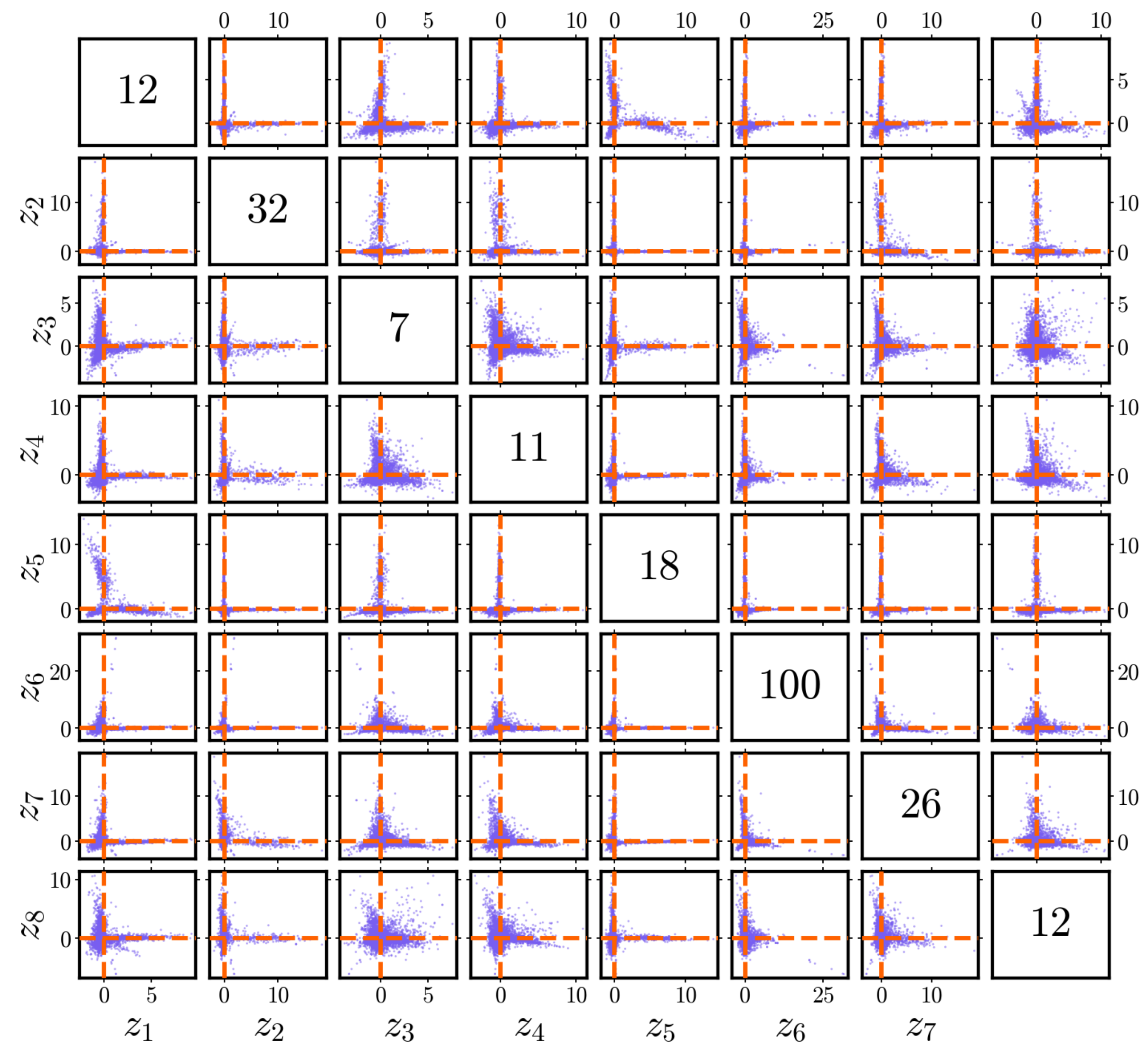
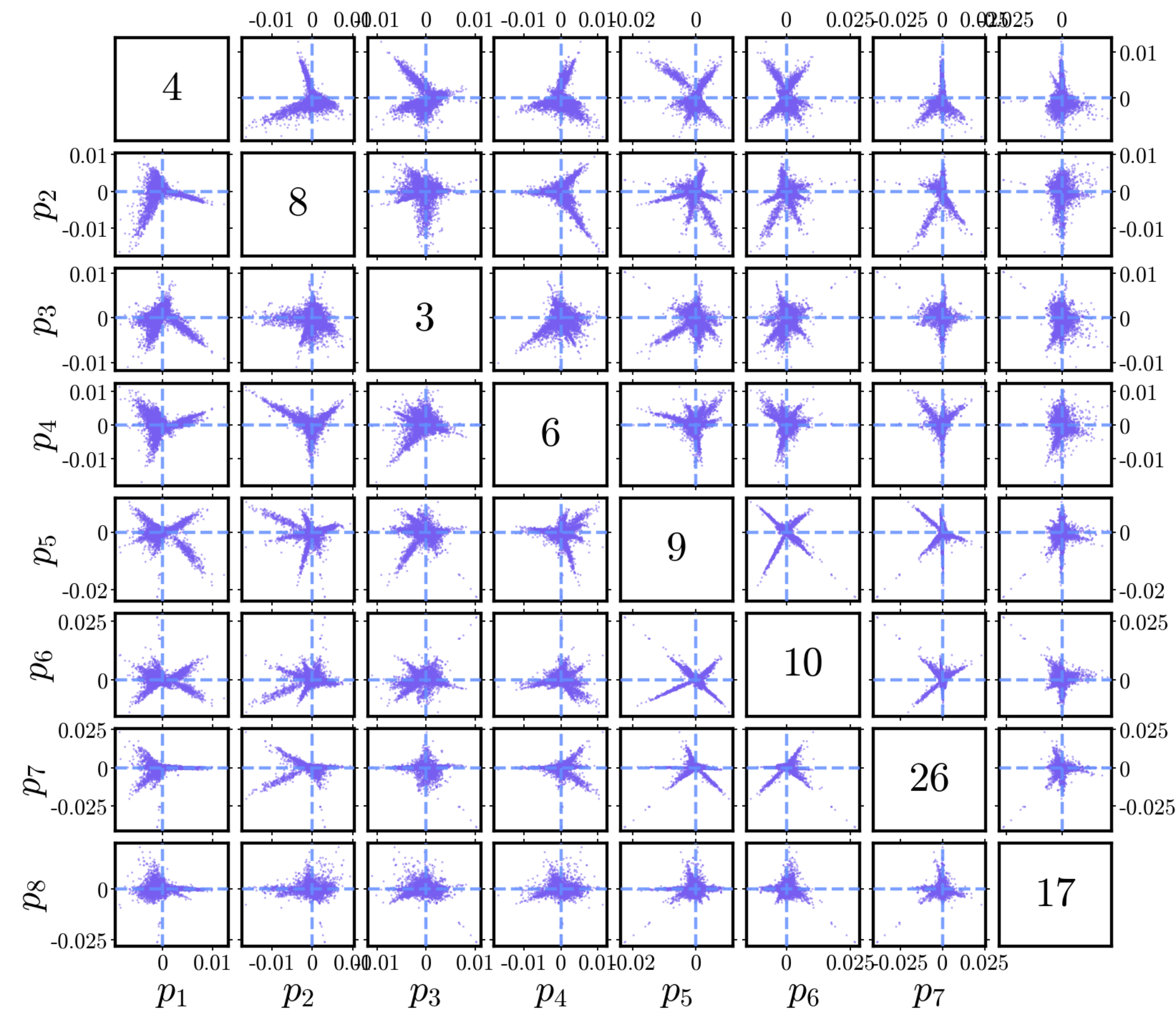


Applications

Given $X \in \mathbb{R}^{d \times n}$, **recover the labels**
 $y \in \{1, \dots, k\}$



NYT bag of words dataset



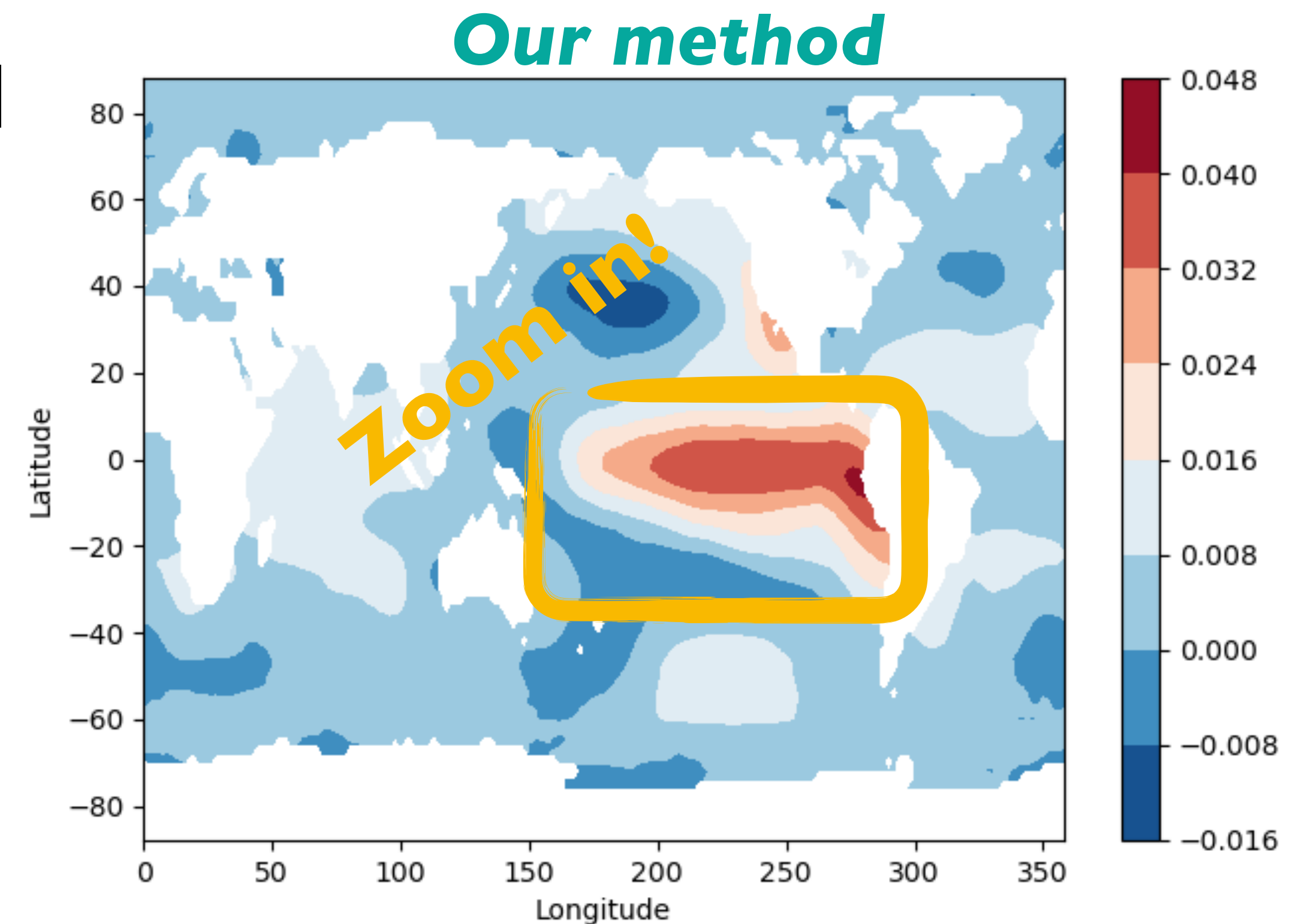
EOF Analysis

Climate researchers want to study **spatial variables** and how they change with time (multidimensional time series data)

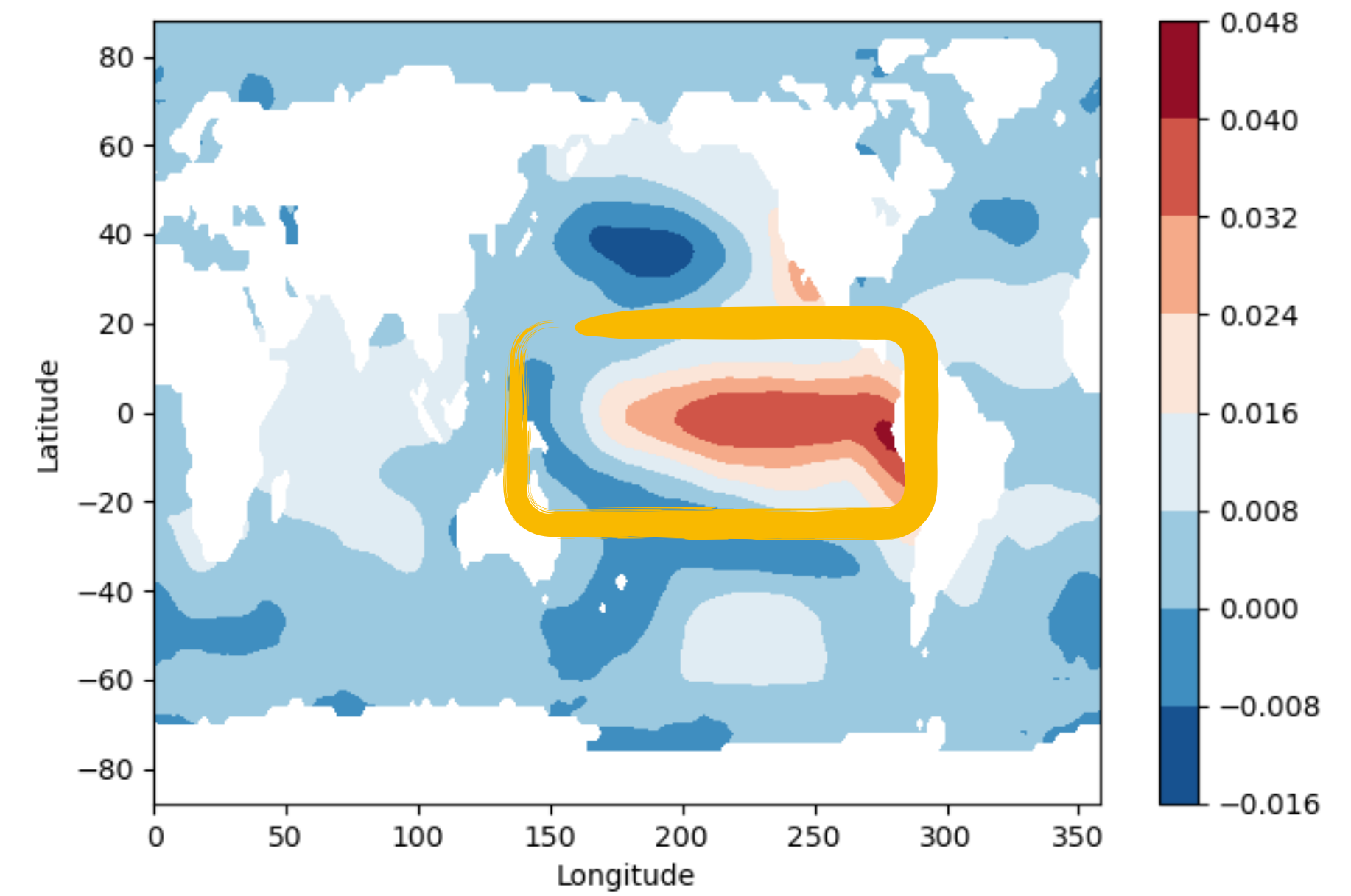
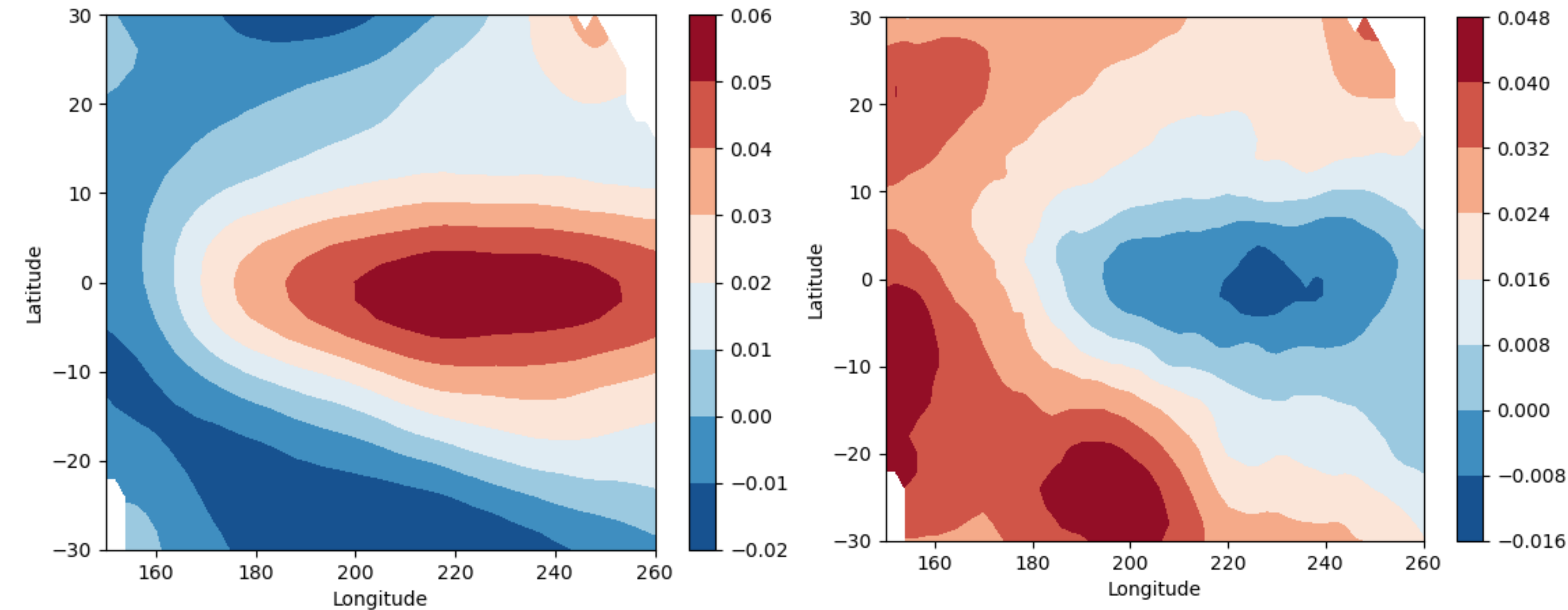
Goal: Fewer **uncorrelated** latent variables (*factors*) of an **anomaly covariance matrix**

Observed field:

$$X = [X_1, \dots, X_d]^\top, \text{ with } X_t = [X_{t_1}, \dots, X_{t_d}] \implies \text{Cov}(X) = \mathbb{E}XX^\top$$



(Real) Real-data applications

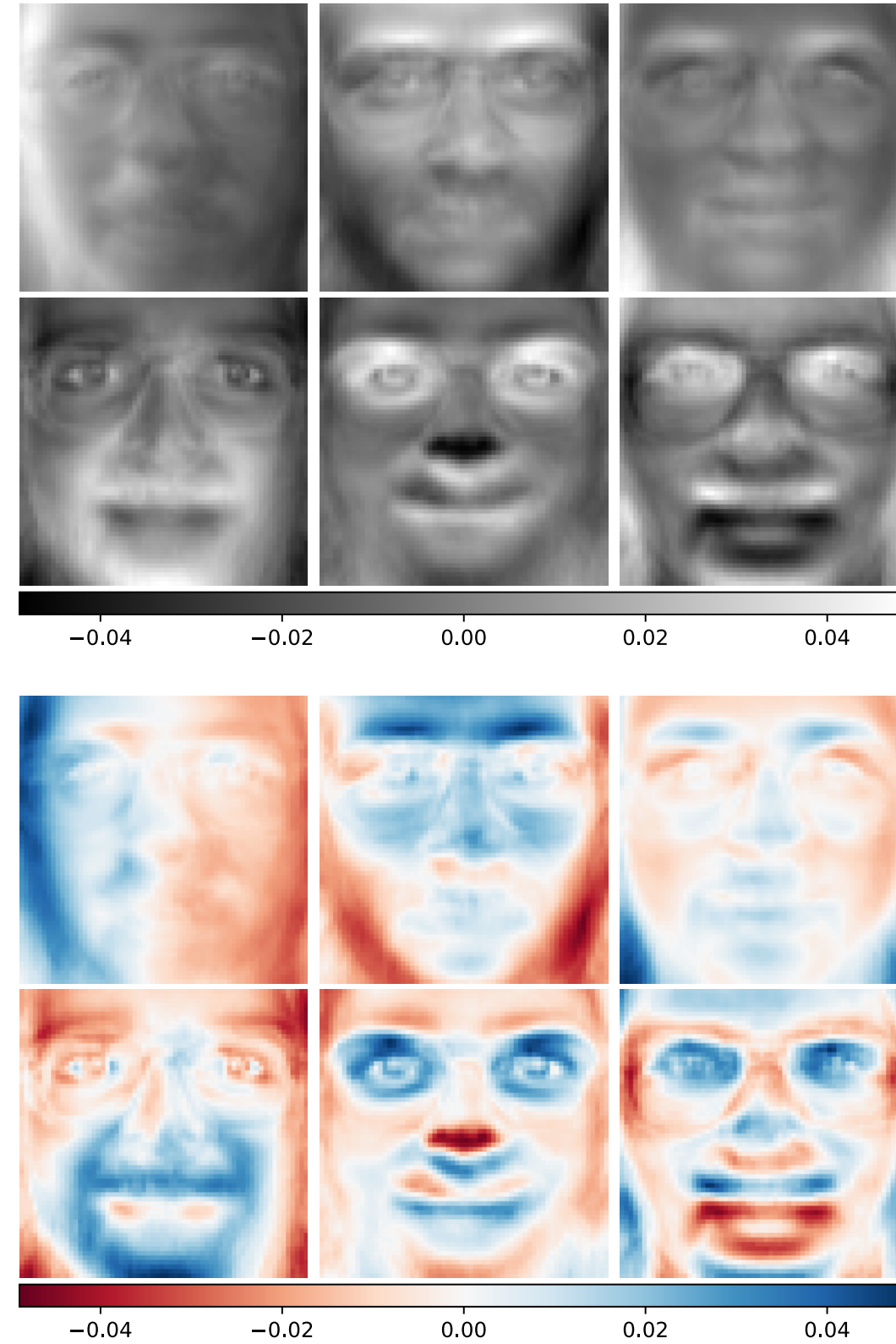


In **sea surface temperature** data, GFA detects **El Niño** warming in the central and eastern Pacific, while orthogonality captures **La Niña**-like cool water patterns.

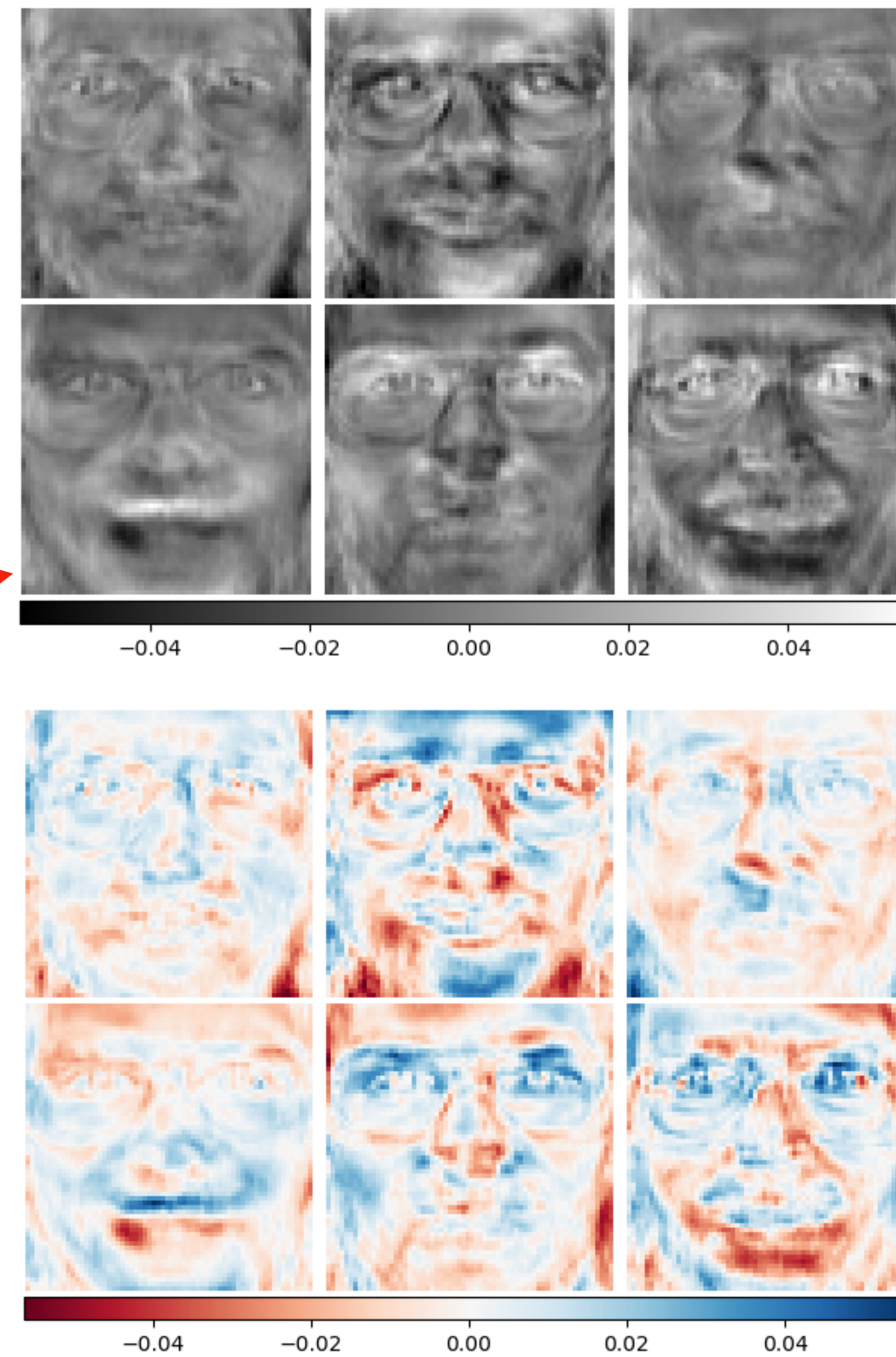


Faces dataset

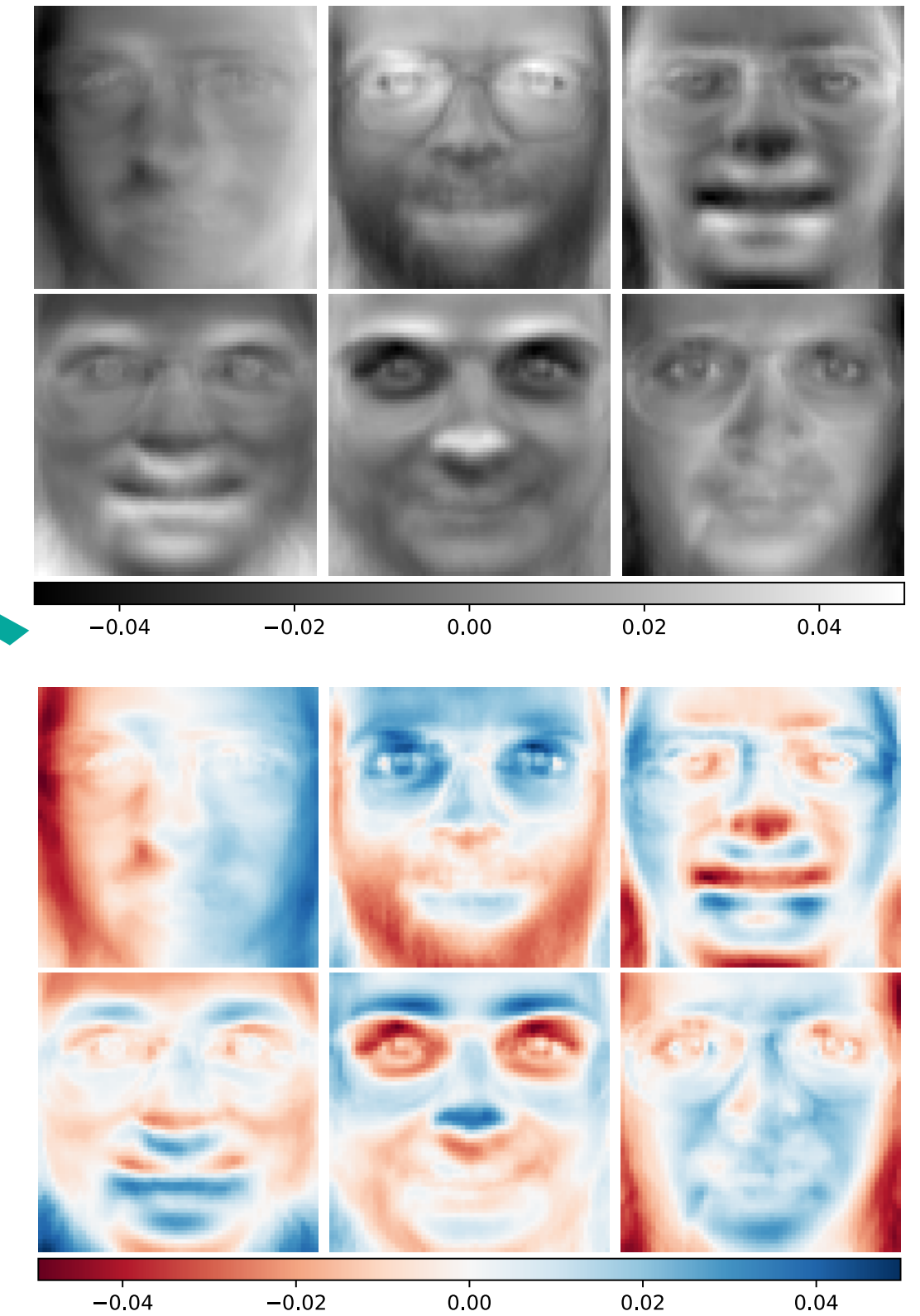
PCA



Varimax



GFA



GFA provides a **more interpretable** factor rotation of high-dimensional datasets **compared to Varimax**

Summary

- ◆ **Constructed a method that uncovers latent (fewer) factors that are a linear combination of correlated variables**
- ◆ **This method is scalable and achieves better statistical performance than other classical methods**
- ◆ **Provided optimization guarantees for this method**
- ◆ **Demonstrate its applicability in real-world datasets**

Gracias!!

